

Chapter 2

Packet Scheduling and Congestion Control

Giovanni Giambene, Mari Carmen Aguayo Torres, Edmundo Monteiro,
Michal Ries, and Vasos Vassiliou

2.1 Introduction

In the framework of *Next-Generation Networks* (NGNs), the IP-based *Radio Access Networks* (RANs) component, though representing one of the most attractive aspects of NGNs, is also the weakest part. The wireless interface is in fact the critical factor in end-to-end *Quality of Service* (QoS) support due in particular to the fluctuations of radio channel conditions and consequent errors. Moreover, the problems arising from the highly varying traffic, wireless transmitter energy consumption (e.g., [ODR05, Ber07]), and the highly desirable user mobility create needs to be addressed before wireless broadband Internet services can be widely and successfully deployed.

Radio Resource Management (RRM) plays a key role in wireless system design. A fundamental element in resource management is *scheduling* that arbitrates among packets that are ready for transmission. Based on the scheduling algorithm, as well as the traffic characteristics of the multiplexed flows, certain QoS levels can be obtained. There are many scheduling algorithms that take care of different parameters, such as deadlines, throughput, channel conditions, energy consumption, and so forth. This chapter will provide an in-depth investigation of several scheduling schemes for wireless systems that are also able to support physical layer adaptivity. In addition to this, *Call Admission Control* (CAC) schemes will also be addressed because they are central elements for resource management with QoS support in wireless systems.

In recent years, cellular and wireless systems have become very popular technologies. On the other hand, communication networks based on *Geostationary Earth Orbit* (GEO) satellites allow the provision of multimedia services covering broad geographic areas. The different *Medium Access Control* (MAC) layer implications for these distinct scenarios are the main subject studied in this chapter. In this chapter, the envisaged technologies are for wireless networks up

G. Giambene (✉)
University of Siena, Italy
e-mail: giambene@unisi.it

to the metropolitan scale: *Universal Mobile Telecommunications System* (UMTS) and *Wideband Code Division Multiple Access* (WCDMA), *Time Division Multiple Access/Time Division Duplexing* (TDMA/TDD), *High Speed Downlink Packet Access* (HSDPA), WiFi, *Orthogonal Frequency Division Multiple Access* (OFDMA) suitable for WiMAX and 3GPP LTE air interfaces, and GEO satellite.

One interesting aspect associated with radio resource management in wireless systems is its impact on higher-layer performance. The particular *Transmission Control Protocol* (TCP) impairments due to buffer congestion phenomena and packet loss events are of special interest here. TCP is today's dominant transport layer protocol for reliable end-to-end data delivery over the Internet. It is designed to use the available bandwidth for the source-destination pair in a fair and efficient way. Its congestion control, originated from wired networks, where congestion is the main cause of packet loss, comes under pressure in wireless networks. Because these networks are characterized by dynamically variable channel conditions, especially due to user mobility, channel fading, and interference conditions, the performance of TCP degrades. The root of this degradation rests in the difficulty for TCP to distinguish between congestion, contention, and channel errors. Moreover, the wireless MAC may cause unfairness for the transport layer congestion control: when more nodes contend for access to the wireless resource, the node that first wins the contention achieves a better capacity (i.e., higher congestion window value). Finally, the standard TCP congestion control mechanism is known to perform poorly over satellite broadband links due to both the large *Round-Trip Time* (RTT) value and the typically high packet error rates.

System efficiency is an important requirement for wireless communication systems to provide broadband services to users. Whereas QoS support is mandatory for end-users who expect a good service level, resource use is a transparent matter to them. System optimization and QoS support are typically conflicting needs that could be solved by means of a suitable *cross-layer system design* aimed at creating and exploiting the interactions between protocols at the logically different *Open System Interconnection* (OSI) reference model architectural layers that otherwise would be operated independently according to the classic OSI layer separation principle. These issues are dealt with in this chapter together with the interaction between resource management and higher (OSI) layers, in particular the transport layer. In this regard, layer 2 choices in traffic management and active queue management schemes are considered.

There is a rich literature dealing with the interactions between RRM and congestion control. In particular, Price and Javidi [Pri04] deal with the interaction between transport layer and MAC layer by performing an integrated rate assignment. Friderikos et al. [Fri04] also study a TCP-related rate adaptation scheme. Hossain and Bhargava [Hos04] analyze the link/physical (PHY) layers influence on TCP behavior. In [Che05], a model is presented for the joint design of congestion control at the transport layer and MAC for *ad hoc* wireless networks. More details are available in [Sun06].

The major results achieved in the COopération Scientifique et Technique (European Cooperation in Science and Technology) (COST) 290 [COS290] that are described in this chapter are summarized as follows: (i) a proposal of a service differentiation scheme for small packets (sensor data and VoIP applications) and an enhanced proposal for QoS support in voice over *Wireless Local Area Network* (WLAN); (ii) the design of resource allocation algorithms for different wireless technologies and different traffic scenarios; (iii) the development of new congestion control protocols suitable for the wireless scenario that use explicit single-bit and multi-bit feedback information; (iv) the selection of MAC layer parameters of IEEE 802.11e to support traffic flows with different requirements in terms of goodput and delay; (v) several cross-layer design proposals involving different layers of the protocol stack to improve the efficiency of the wireless interface.

This chapter is organized as follows. After this introduction, service differentiation is described, then RRM schemes for wireless systems are addressed, followed by considerations on the impact on the transport layer performance of the wireless scenario and RRM. The substantial cross-layer issues are then presented, focusing on different possible interactions.

2.2 Service Differentiation

Service differentiation is a key issue for the support of new applications with demanding QoS requirements in the future Internet. This section addresses service differentiation mechanisms in wired and wireless networks. Section 2.2.1 provides a survey on scheduling disciplines used for service differentiation. The remaining sections summarize current proposals for service differentiation of noncongestive applications (described in Section 2.2.2) and for QoS support in voice over WLAN (presented in Section 2.2.3).

2.2.1 Scheduling for Service Differentiation

Scheduling in wireline communication systems has a long research tradition. An introduction to this wide area can be found in [Zha95, Nec06]. Coming from the wireline domain, certain characteristics of wireless links make it difficult to apply directly many existing scheduling algorithms. One such characteristic is the relatively high loss probability of data frames due to transmission errors. Most wireless systems therefore apply *Automatic Repeat reQuest* (ARQ) mechanisms to recover from losses [Com84]. The capacity needed by retransmissions introduces additional delay and consumes resources on the transmission link not foreseen by conventional scheduling algorithms. A second and much more severe problem is the fact that many state-of-the-art wireless systems like HSDPA employ adaptive modulation and coding, leading to a variable data rate over time toward each user. This derives from the time-varying

nature of the radio channel and makes it difficult to specify directly a capacity of the shared link, as it would depend on which users are being served in a particular scheduling round.

Even though a time-variant radio channel may impose problems to conventional scheduling schemes, it gives rise to *opportunistic* schedulers, also known as *channel-aware* schedulers, which favor terminals with temporarily good channel conditions. Knowledge and exploitation of channel conditions to various users can significantly increase system capacity. One popular approach is the *Proportional Fair* (PF) scheduler, which bases its scheduling decisions on both the instantaneous and the average channel quality to particular terminals. The performance of the PF scheduler in an HSDPA environment under particular consideration of imperfect channel knowledge has been investigated in [Kol03].

Classic channel-aware schedulers do not include mechanisms to provide QoS guarantees. Several approaches exist to combine channel-aware schedulers with traditional QoS-aware scheduling approaches that are based on service differentiation.¹ In the following, a classification and summary of popular approaches is presented.

2.2.1.1 Scheduling in Wireline Networks

This part summarizes the major service differentiation approaches for wireline networks. Further approaches exist (e.g., rate-controlled approaches), but will not be discussed here.

Static Prioritization: A *Static Prioritization* (SP) of higher-priority traffic is the simplest form of service differentiation. In this approach, the high-priority traffic always gets the best possible service quality. On the other hand, this may lead to excessive interscheduling gaps and even the starvation of lower-priority traffic. Moreover, it does not allow for exploitation of the delay flexibility inherent in certain high-priority traffic classes. In particular, the use of this scheme in a wireless environment partially prevents an efficient utilization of air interface resources, as is possible in opportunistic scheduling schemes.

Deadline-based schemes: The most basic deadline-based scheme is the *Earliest Due Date* (EDD) algorithm, proposed in [Jac55, Liu73], which assigns each packet a delivery deadline based on its flow's QoS parameters. This allows a more flexible utilization of the packet delay tolerance by preferring packets from lower-priority traffic if high-priority traffic is not in danger of violating its deadline. This approach enables a more efficient use of air interface resources in combination with opportunistic scheduling mechanisms.

Fair queuing: Fair queuing typically goes back to the *General Processor Sharing* (GPS) model, which bases the link sharing on a fluid-flow model with an infinitesimally small scheduling granularity. This makes it a theoretical

¹ This part was published in *International Journal of Electronics and Communications* (AEÜ), Vol. 60, No. 2, M. C. Necker, A comparison of scheduling mechanisms for service class differentiation in HSDPA networks. Copyright © 2006 Elsevier.

model, and many practically realizable approximations have been presented in the literature, such as the well-known *Weighted Fair Queuing* (WFQ) algorithm [Dem89], also known as *Packet-based GPS* (PGPS) [Par93]. GPS-based approaches are very popular and widely deployed in wireline networks. In wireless systems, GPS-based approaches especially suffer from the problems mentioned before. The problem of ARQ-mechanisms has been studied for example in [Kim05]. Moreover, in HSDPA and in many other wireless systems, variable and individual data rates toward each terminal violate a basic assumption of GPS-based approaches.²

2.2.1.2 Scheduling in Wireless Networks

The major service differentiation approaches for wireless networks are described in the following paragraphs.

Maximum Carrier to Interference (C/I) ratio: The Maximum C/I scheduler, also known as *Signal-to-Noise Ratio* (SNR)-based scheduler, is the most simple channel-aware scheduler. It bases its scheduling decision on the absolute instantaneous channel quality reported by each *User Equipment* (UE) in each scheduling round. Consequently, it maximizes the overall system capacity and the aggregate throughput. The main disadvantage of this approach is the inherent unfairness, especially on small timescales. Because of its absolute metric, it may cause excessive scheduling gaps for, and even starvation of, users in unfavorable positions. This will lead to excessive cross-layer interactions with higher-layer protocols. Consequently, the Maximum C/I scheduler is only of theoretical interest, for example as a performance reference or benchmark.

Proportional Fair (PF): The PF scheduler [Jal00] overcomes the fairness problems of the Maximum C/I approach by basing its scheduling decisions on the ratio between the currently achievable data rate, $R_k(t)$, and the averaged data rate over the recent past, $\bar{R}_k(t)$, to a particular terminal, that is,

$$tag_k = \frac{R_k(t)}{\bar{R}_k(t)}, \quad (2.1)$$

for flow k at time instant t . $\bar{R}_k(t)$ is updated as follows:

$$\bar{R}_k(t) = \frac{1}{\tau} R_k(t) + \left(1 - \frac{1}{\tau}\right) \bar{R}_k(t - T_{TTI}), \quad (2.2)$$

² This part was published in *International Journal of Electronics and Communications* (AEÜ), Vol. 60, No. 2, M. C. Necker, A comparison of scheduling mechanisms for service class differentiation in HSDPA networks. Copyright © 2006 Elsevier.

where T_{TTI} is the length of a scheduling interval, that is, the length of a *Transmission Time Interval (TTI)* in HSDPA. The weighting factor τ is a time constant.

The above equation immediately raises the question of how to update $\bar{R}_k(t)$ when the buffer of the user is empty. Different possibilities exist, which have been studied in detail in [Fel06].

In order to add QoS capabilities to an opportunistic scheduling scheme, it needs to be combined with class differentiation mechanisms. The following paragraphs discuss several options.³

Hierarchical scheduling: In hierarchical scheduling, there is a double hierarchy. One scheduler determines, for each active traffic class, which flow currently has the highest priority. The other, an interclass scheduler, also known as a link-sharing scheduler, decides on the traffic class to be served in the next scheduling round. As typical examples, here a PF scheduler is considered as a traffic class scheduler, and both an SP and a *Weighted Round-Robin (WRR)* scheduler as interclass schedulers.

Deadline-based schemes: Deadline-based schemes extend an opportunistic scheduler with a delay-dependent component. Each traffic class has its own maximum delay T_k taken into account by the scheduler. Two examples of this approach will be considered here, namely the *Channel-Dependent Earliest Due Date (CD-EDD)* and the *Exponential Rule (ER)* algorithm.

CD-EDD scheduling has been proposed in [Kha04]. It is a combination of the PF approach and an EDD component according to the following formula:

$$tag_k = w_k \underbrace{\frac{R_k(t)}{\bar{R}_k(t)}}_{\text{PF term}} \underbrace{\frac{W_k(t)}{d_k(t)}}_{\text{EDD term}} = w_k \frac{R_k(t)}{\bar{R}_k(t)} \frac{W_k(t)}{T_k - W_k(t)}, \quad (2.3)$$

where, for the k -th flow, w_k is the k -th weighting factor, which here is set to 1; $W_k(t)$ denotes the waiting time of the *Head of Line (HOL)* packet in the queue of the k -th flow, T_k is the maximum allowable delay of a packet in the k -th flow, and $d_k(t) = T_k - W_k(t)$ is the time remaining to the deadline.

The graph of the EDD term is shown in Fig. 2.1. The exponential rise (toward infinity) as the delay of the HOL-packet, W_k , tends toward T_k , forces the EDD term to quickly dominate the scheduling tag.

For low HOL-packet delays, the EDD term, exponentially tending to zero, gives the flow a rather low priority. A modification that ameliorates this low-priority exponential effect is to let the PF algorithm do the scheduling as long as no HOL packet is in danger of violating its deadline [Bar02].

³ This part was published in *International Journal of Electronics and Communications (AEÜ)*, Vol. 60, No. 2, M. C. Necker, A comparison of scheduling mechanisms for service class differentiation in HSDPA networks. Copyright © 2006 Elsevier.

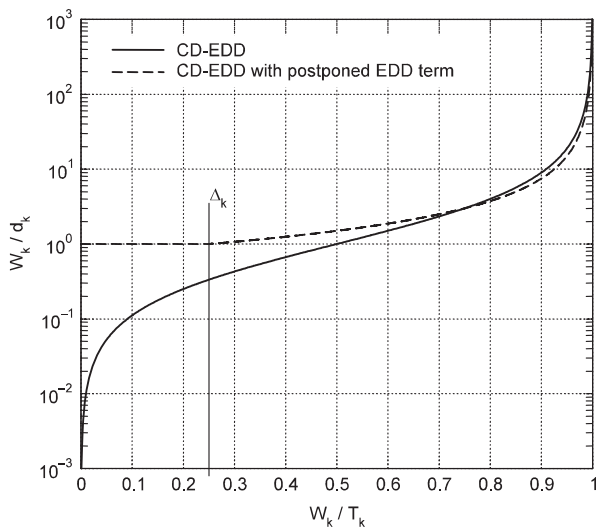


Fig. 2.1 EDD weighting function in Equations (2.3) and (2.4)

Defining a minimum delay Δ_k for each flow k , the scheduling tag can be redefined as:

$$tag_k = a_k \frac{R_k(t)}{\bar{R}_k(t)} \left(\frac{\max(0, W_k(t) - \Delta_k)}{T_k - W_k} + 1 \right), \quad (2.4)$$

where the term within brackets (i.e., postponed EDD term) is plotted in Fig. 2.1.⁴

Exponential Rule (ER) scheduling: As a second example of combining the PF approach with delay sensitivity, the ER scheduling approach will be considered [Sha01]. The idea behind ER is to rely on the PF algorithm for regular situations and to equalize the weighted delays of the queues of all flows if their differences are large. This makes it very similar to the above-described modification of the CD-EDD algorithm. The scheduling tags are calculated as follows:

$$tag_k = a_k \frac{R_k(t)}{\bar{R}_k(t)} \exp \left(\frac{a_k W_k(t) - \overline{aW}(t)}{1 + \sqrt{\overline{aW}(t)}} \right), \quad a_k = -\frac{\log(\delta_k)}{T_k}, \quad (2.5)$$

where T_k is again the maximum allowable delay, W_k is the delay of the HOL packet, and δ_k is the largest probability with which the scheduler may violate the delay deadline. $\overline{aW}(t)$ is defined as $\frac{1}{K} \sum_k a_k W_k(t)$, where K denotes the total number of flows.

⁴ This part was published in *International Journal of Electronics and Communications (AEÜ)*, Vol. 60, No. 2, M. C. Necker, A comparison of scheduling mechanisms for service class differentiation in HSDPA networks. Copyright © 2006 Elsevier.

Modified Largest Weighted Delay First (M-LWDF): This scheduler was proposed in [And02] and proved to be throughput-optimal in the sense that it can handle any offered traffic.

For a performance evaluation of these scheduling algorithms in multiservice HSDPA scenarios, interested readers are referred to [TD(06)015]; see also Section 3.2.1.1.

The following sections describe current research work on service differentiation for wireless and wired environments.

2.2.2 Service Differentiation for Noncongestive Applications

This section deals with a service differentiation scheme for small packets, particularly suitable for sensor and VoIP applications. It derives from a new service strategy, called *Less Impact Better Service* (LIBS), according to which “noncongestive” traffic (i.e., small packets at low rates that require minor service delays and hence cause minor queuing delays) gets some limited priority over long packets. The limitation is strictly associated with the cumulative service impact of this prioritization on long packets.

2.2.2.1 Noncongestive Queuing

In [TD(05)013, Mam05, Mam07], the authors have shown that, based on service thresholds, service differentiation can be achieved for noncongestive applications, such as sensor applications or other types of applications that use small packets and rates with almost zero cost on congestive applications.

Typical service paradigms assume resource demand exceeding resource supply, thus focusing on bandwidth sharing among flows. Other service paradigms incorporate a proportional service scheme, where bandwidth allocation is made in proportion to the demand. However, both perspectives lack a delay-oriented service discipline. Management of delay-oriented services has been traditionally based on delay requirements of some applications, which are eventually reflected in the prioritization during scheduling. Thus, service disciplines are primarily application-oriented and have the inherent property to better satisfy some applications more, rather than the satisfying of more applications. With a goal of satisfying more users, a system-oriented service discipline is considered in the following. This service approach promotes, and thereby fosters, “noncongestive” traffic. To avoid starvation and also significant delay impact on congestive traffic, noncongestive traffic prioritization is confined by corresponding service thresholds. From a user perspective, the key QoS requirements of applications that use small data packets and rates (and are also intolerant to long delays) are satisfied, whereas other applications suffer almost zero extra delays. This

service differentiation scheme, called *Noncongestive Queuing* (NCQ) [TD(05)013, Mam05, Mam07], enables the separation between noncongestive flows due to real-time applications and other flows that use small packets as well.

The key idea of NCQ derives from the operational dynamics of gateways: they may service small packets instantly. Noncongestive packets do not cause significant delays and hence should not suffer from delays. Although this approach here sounds straightforward, the system properties and design details reveal interesting dynamics. The simplicity of NCQ's core algorithm reduces implementation and deployment efforts. NCQ does not require any modification at the transport layer or packet marking; a minor modification of the gateway software is sufficient.

2.2.2.2 Performance Evaluation

In order to demonstrate the potential of NCQ, a simple *ns-2* (*Network Simulator*, version 2 [ns207]) based experiment was carried out, comparing NCQ with DropTail queuing (which drops incoming packets when the queue is full). The simple *dumbbell topology* (two sources and two sinks interconnected through a bottleneck link) was used with the aim to measure *goodput* for both congestive (e.g., FTP) and noncongestive applications (e.g., sensor data).

As may be seen in Fig. 2.2(a), noncongestive flows achieve significant performance gains (e.g., 4.9 times, in case of 70 flows) in terms of goodput. Although noncongestive traffic is clearly favored by NCQ, occasionally better performance for the congestive flows in Fig. 2.2(b) may be observed (only minor differences). This is not unreasonable: the impact of timeouts caused by short packets is more significant for noncongestive flows compared with that for long packets.

The proposed service paradigm has an impact on other performance measures as well, such as energy expenditure. This is a significant issue for energy-limited devices (e.g., sensors or VoIP mobile devices). The energy savings are achieved through reduction of the transmit communication time. Savings vary depending on the device itself, the communication pattern, the network contention, and so forth. In [Mam05], the authors show that NCQ improves energy efficiency and real-time communication capability of sensor devices and applications, respectively, without causing any *goodput* losses to congestive flows. In [TD(05)013, Mam05], taking an analytic approach to the NCQ mechanism, the authors experimented with different traffic thresholds and traffic class proportions to demonstrate the overall system behavior when NCQ parameters change. Outcomes of the mentioned studies have proved the usefulness of NCQ as a candidate for service differentiation for noncongestive applications.

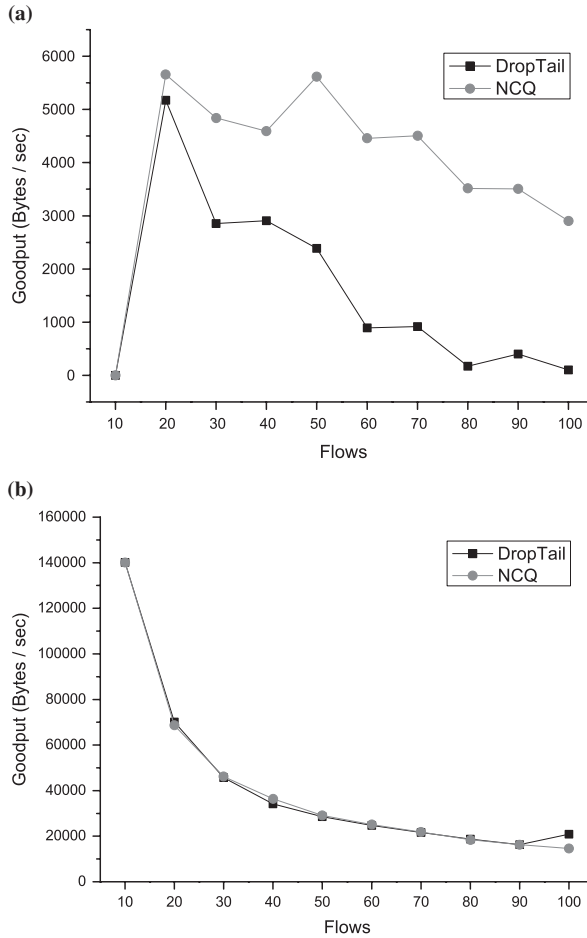


Fig. 2.2 Average *goodput* of (a) noncongestive flows and (b) congestive flows

2.2.3 *QoS for Voice Communications Over WLANs*⁵

Voice communications are one of the leading applications that benefit from the mobility and increasing bit rates provided by current and emerging WLAN technologies. *Voice over WLAN* (VoWLAN) is one important application for WLANs.

Nowadays, users expect toll-quality service regardless of the medium (wired vs. wireless) and the switching technology employed. Multiservice networks require special prioritization protocols to ensure good voice performance. The

⁵ Parts of this Section were published in [Vil06b]. Copyright © 2006 LNCS-Springer.

IEEE 802.11e standard [IEEE05] defines mechanisms to meet the QoS requirements of various applications, such as voice and video services (see Section 2.5.1 for more details). In the near future, it is expected that the IEEE 802.11e interface cards will take over the WLAN market, replacing the use of legacy IEEE 802.11 [IEEE99] interface cards in most WLAN applications, though complete migration will require several years, given the wide-scale use of legacy IEEE 802.11 in the marketplace today. Hence, the number of networking scenarios where legacy IEEE 802.11-based stations and IEEE 802.11e-based stations will coexist and interoperate for a period will likely be significant.

However, it is observed that the *Enhanced Distributed Channel Access*, EDCA in IEEE 802.11e (see the following section), performs poorly as the network load increases, mainly due to the higher probability of collision. This reason has led many researchers to design techniques aiming to improve the EDCA performance [Kwo04, Ma04]. The two main drawbacks of proposals to date are (i) their implementation requires important modifications to the IEEE 802.11e specifications; and (ii) their inability to meet QoS requirements for the multimedia applications in the presence of legacy *Distributed Coordination Function* (DCF)-based stations. In the following, how these two drawbacks are addressed by introducing a novel IEEE 802.11e-compliant mechanism is set out. This mechanism is capable of providing QoS guarantees to voice services even under scenarios where legacy DCF-based stations are present. The main objective has been to design a scheme compatible with the IEEE 802.11 standards, including DCF and EDCA mechanisms. Simulation results show that the new scheme outperforms EDCA.

2.2.3.1 B-EDCA: A New IEEE 802.11e-Based QoS Mechanism

Because of the vast legacy IEEE 802.11 infrastructure already in place, IEEE 802.11e-based systems will be required to properly interoperate with the existing mechanisms, such as DCF. Under such scenarios, EDCA has been shown to be unable to meet the QoS of time-constrained services, in particular voice communications. Based on these limitations and under the constraint of ensuring compatibility with the existing mechanism as a key element for its successful deployment, a new IEEE 802.11e-based QoS mechanism capable of providing QoS support to time-constrained applications has been introduced in [Vil06a, TD(06)038].

Bearing in mind that DCF and EDCA mechanisms may have to interwork, the standards committee has set up the system parameters given in Table 2.1. These values have been identified in order to ensure compatibility between both services, with the EDCA mechanism being able to provide QoS guarantees to time-constrained applications, namely voice and video traffic. As shown in Table 2.1, EDCA makes use of a shorter contention window for voice and video applications.

To introduce the proposal here, a closer look at the mode of operation of DCF and EDCA schemes is needed, particularly on the role played by the IFS (or AIFS; *Arbitration Inter-Frame Space*) parameter. The IFS (AIFS) interval is used in the Idle state: when the station becomes active, it has to sense the

Table 2.1 Parameter settings specified in the WiFi standard [IEEE05, IEEE99, Vll06b]; in particular, *Contention Window* (CW) and *Inter-Frame Space* (IFS). Four access classes are considered in EDCA, such as *voice* (Vo), *video* (Vi), *best-effort* (Be), and *background* (Bk) (copyright © 2006 LNCS-Springer [Vll06b])

	AC	IFS	CW _{min}	CW _{max}
DCF	–	$2 \times \text{Slot_time} + \text{SIFS}$	31	1023
EDCA	Vo	$2 \times \text{Slot_time} + \text{SIFS}$	7	15
	Vi	$2 \times \text{Slot_time} + \text{SIFS}$	15	31
	Be	$3 \times \text{Slot_time} + \text{SIFS}$	31	1023
	Bk	$7 \times \text{Slot_time} + \text{SIFS}$	31	1023

channel during an interval whose length is determined by IFS; if the channel sensed is free, the station can initiate the packet transmission. Otherwise, the station executes the backoff algorithm.

According to the current DCF and EDCA standards, the same values for the IFS parameter should be used regardless of the state in which the station is (see Table 2.1). Based on the previous observation, a different set of IFS values is proposed for use, depending on the state in which the station is. The *Hybrid Coordination Function* (HCF) operation, however, cannot be compromised, and in particular, the technique must ensure that it holds the highest priority at all times.

It is possible to introduce the following definitions:

- *Defer state*: this state is entered when the station wants to transmit a frame, but the medium is busy.
- *Backoff state*: this state is entered after a Defer state while the station waits a random number of slots before transmitting to avoid collisions.

In every transfer from a Defer state to a Backoff state, a different parameter is proposed here for use. Denoted here as BIFS, it is equivalent to the IFS. It is proposed to set its value to one slot time, that is, $\text{BIFS} = 1$, for voice and video services. In this way, the performance of voice and video applications improved considerably, and likewise their priorities with respect to other flows (included the traffic generated by DCF-based stations) increased. This setting also ensures that the *Hybrid Coordinator* (HC) will keep the highest priority. According to this mechanism, the stations must wait at least one additional *Short Inter-Frame Space* (SIFS) time ($\text{AIFS}[AC] = 2 \times \text{SIFS} + \text{aSlotTime}$) (AC stands for Access Class) only when the backoff counter is equal to zero. In turn, the HC is allowed to take the control at the end of the IFS. The use of the set of values for BIFS to 1-1-3-7 for voice, video, best-effort, and background traffic flows, respectively, is proposed.

In Fig. 2.3, the instances where the BIFS parameter should be used are explicitly indicated. This is essentially the major change with respect to the current EDCA standard. The waiting time required to continue decrementing the backoff counter used by the time-constrained applications is effectively reduced to the minimum acceptable value by means of the B-EDCA proposal.

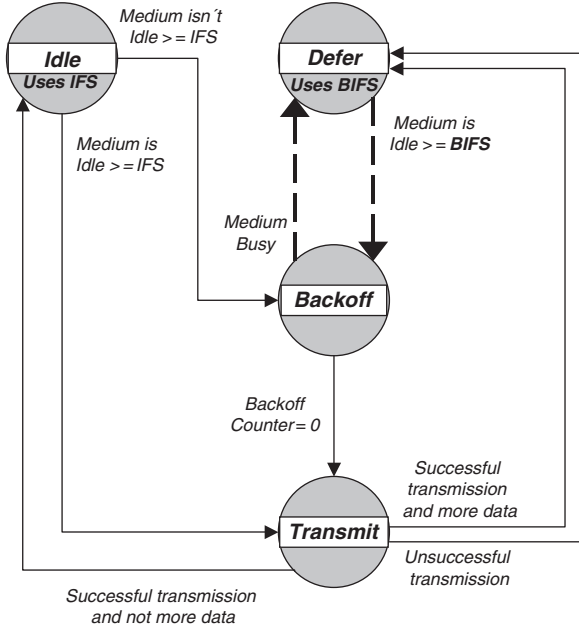


Fig. 2.3 The B-EDCA proposed mechanism [Vil06b] (copyright © 2006 LNCS-Springer [Vil06b])

This value is fully compatible with the operation modes of DCF and *HCF Controlled Channel Access* (HCCA) functions.

2.2.3.2 Performance Evaluation

In this section, a performance analysis is carried out in order to assess the effectiveness of the proposed mechanism. For the following study, the OPNET Modeler tool 11.0 is used [OPNET04], which already supports the IEEE 802.11 DCF simulator. Both EDCA and B-EDCA mechanisms have been integrated into the simulator for this study.

In the simulations, an IEEE 802.11a WLAN is modeled, consisting of several wireless stations and an *Access Point* (AP) that also serves as a sink for the flows coming from the wireless stations. The use of three different types of wireless stations is considered: DCF-compliant stations and EDCA and B-EDCA QoS-aware stations. EDCA- and B-EDCA-based stations support four different types of services, such as Vo, Vi, Be, and Bk, as defined in Table 2.1. This classification is in line with the IEEE802.1D standard specifications.

Figure 2.4 shows the voice, video, DCF, and Global normalized throughput obtained when using EDCA and B-EDCA methods. In the case of the voice service, the proposed B-EDCA scheme is able to provide better QoS guarantees than is EDCA. Taking into account that the maximum acceptable loss rate for the voice service is 5%, it is clear from the results that EDCA is unable to

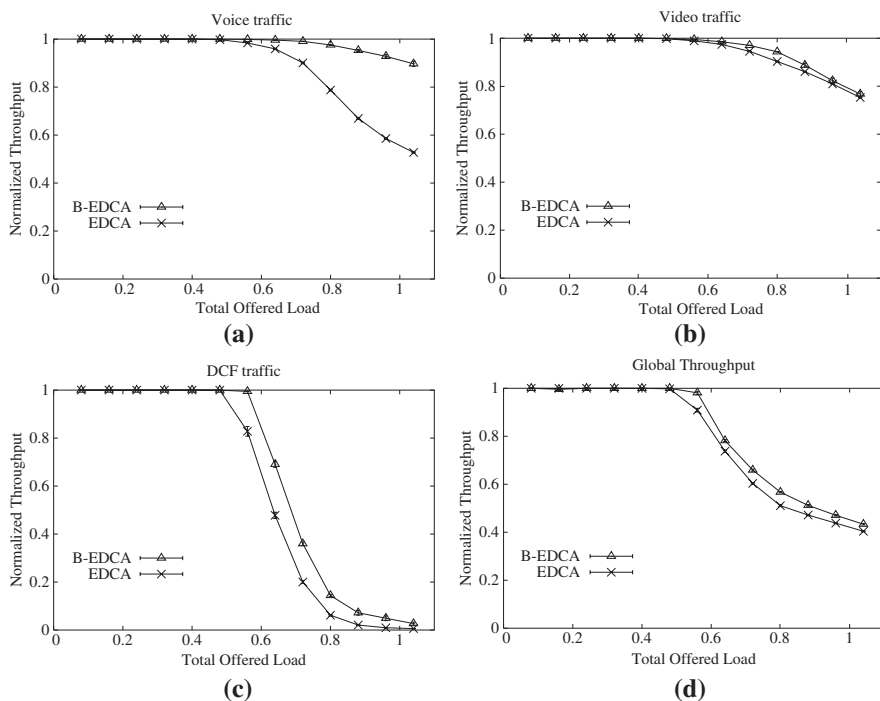


Fig. 2.4 Average normalized throughput: (a) voice, (b) video, (c) DCF traffic, and (d) total traffic (copyright © 2006 LNCS-Springer [Vil06b])

provide such guarantees for load exceeding 65% of the network nominal rate; whereas B-EDCA is able to provide such guarantees for load up to 90%. Figures 2.4(b)–(d) show that the new scheme does not penalize the rest of the traffic. In fact, it is able to improve slightly its throughput. Figure 2.4(d) shows the overall throughput for all the services under study. It is clear that B-EDCA exhibits the highest normalized throughput. This is due to the reduction of the collision rate with respect to EDCA.

The above results show that B-EDCA outperforms EDCA providing better QoS guarantees not only for the voice service but also for the video one. B-EDCA is able to reduce the number of collisions encountered by the voice traffic by one half with respect to the EDCA mechanism. It has also been shown in [Vil06b] that B-EDCA outperforms EDCA under different traffic scenarios, providing better QoS guarantees to the voice service.

2.3 Radio Resource Management Aspects

The current increase of multimedia services brings a new challenge for developing new RRM algorithms. Because of service and traffic differentiations, there is the need for research in the field of RRM. This section deals with

selected RRM aspects concerning different wireless communication technologies, such as WCDMA, HSDPA, IEEE 802.16d/e, IEEE 802.11, and an OFDMA-based air interface.

In the first three sections, novel RRM algorithms are described for IMT-2000 systems (i.e., *third-generation cellular* [3G] systems). Then, a very promising approach for OFDMA networks is described in the following section that allows for effective interference coordination of neighboring base stations. In the final two sections, content-based resource allocation mechanisms are introduced for WLAN networks.

2.3.1 Orthogonal Variable Spreading Factor Code Allocation Strategy Using Genetic Algorithms

Orthogonal Variable Spreading Factor (OVSF) codes have been proposed for the data channelization in WCDMA access technology of IMT-2000. The OVSF codes are the resources that should be commonly used by all system subscribers. Some allocation/reallocation methods for OVSF codes were already proposed and used in [Tse01, TD(05)011]. The common purpose of all of them is to minimize the blocking probability and the reallocation codes cost so that more new arriving call requests can be supported. Efficient channelization code management results in high code utilization and increased system capacity. The probability of code blocking due to the inappropriate resource allocation will be thus minimized.

A *Genetic Algorithm* (GA) is a computational model inspired by evolutionary theory. The algorithm encodes a potential solution to a specific problem on a simple chromosome-like data structure and applies recombination to these structures to preserve the good information. GAs are often viewed as optimizers; the range of problems where GAs have been applied is quite broad [Bak85]. Figure 2.5 describes how a GA operates. An implementation of a GA begins with a population of (typically random) chromosomes. In the next step, one evaluates these structures and allocates reproductive opportunities in such a way that those chromosomes that represent a better solution to the target problem are given more chance to “reproduce” than are those chromosomes that are poorer solutions. Running the GA, the selected individuals will be recombined in order to generate the next population of chromosomes (i.e., the next generation).

The problem of using a GA as a method for allocation/reallocation of OVSF codes is investigated in [Min00]: each chromosome in a population represents an OVSF tree structure (each chromosome is obtained/encoded by reading the codes by levels in the OVSF tree from the root to the last leaf); the most adapted chromosome that will solve the allocation code request is searched. OVSF tree structures will be searched that have the orthogonality condition fulfilled for all active codes. In this technique, the incoming code rate requests are generated

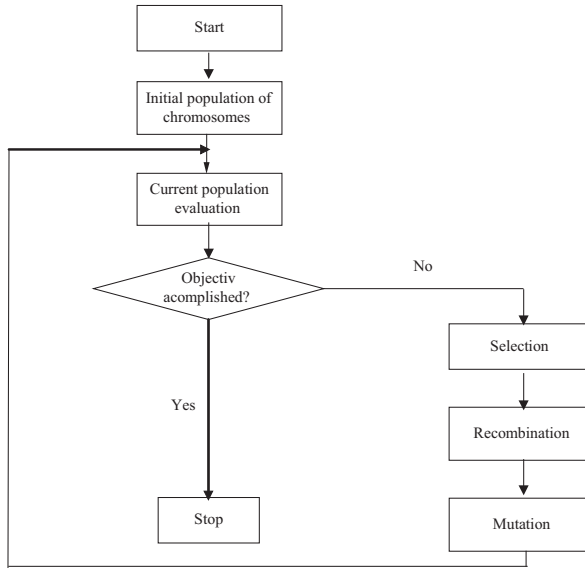


Fig. 2.5 Genetic algorithm diagram

onto the OVSF partially occupied tree, and each request is managed by a GA. At the end of the GA, a new structure of the OVSF tree having solved the code rate request is proposed.

According to the results shown in [Bak85, TD(05)011], GA-based strategies for allocation/reallocation of OVSF codes in WCDMA achieve (with a reasonable computation load in terms of need generations) a blocking probability performance comparable with that obtained by using deterministic and computationally-complex allocation methods.

2.3.2 Novel Buffer Management Scheme for Multimedia Traffic in HSDPA

This section is concerned with the management of multimedia traffic over HSDPA. The current HSDPA architecture provides opportunity to apply buffer management schemes to improve traffic QoS performance and resource use. A novel buffer management scheme [Beg06] based on priority queuing, the *Time-Space-Priority* (TSP) scheme, is proposed for QoS management of single-user downlink multimedia traffic with diverse flows in HSDPA Node-B.

TSP is a hybrid priority queuing scheme that combines time priority and space priority with a threshold to control the QoS parameters (loss, delay, and jitter) of concurrent diverse flows within a multimedia stream. *Real-Time* (RT) delay-sensitive flows, such as video or voice packets, are queued in front of *Non-Real-Time* (NRT) flows, such as e-mail, *short message service* (sms) or file

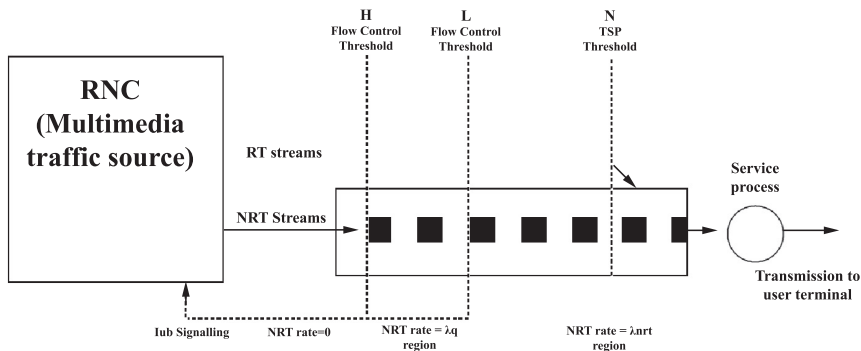


Fig. 2.6 Extended TSP scheme with dual-threshold rate control for loss-sensitive flows in the multimedia traffic

downloads, to receive non-preemptive priority scheduling for transmission on the shared channel. NRT flows are granted buffer space priority to minimize loss by restricting RT flow admission into the buffer. The extended TSP scheme includes thresholds for flow control applicable only to delay-insensitive NRT flows.

As illustrated in Fig. 2.6, TSP aims at optimizing the conflicting QoS requirements of each flow in the multimedia stream by means of a threshold, N , which restricts RT packets (or data units) admission into the shared buffer queue thereby ensuring space priority for the NRT flow to minimize NRT loss. At the same time, restricting RT admission with threshold N reduces RT jitter. In order to ensure minimum RT transmission delay, RT data units are queued in front of the NRT data units on arrival in order to receive priority transmission on the shared channel (i.e., time priority). In addition to QoS optimization [TD(05)048], TSP also provides an efficient way to utilize transmission buffers. Further details on the HSDPA air interface can be found in Section 3.2.1.1.

2.3.3 Power Control for Multimedia Broadcast Multicast Services

The importance of resource-use efficiency for multicasting in UMTS (i.e., *Multimedia Broadcast Multicast Service*; MBMS) has been presented in [Bar03]. In 3G systems, which are based on the CDMA technique where all users can share a common frequency band, interference control is a crucial issue. In WCDMA, a group of *power control functions* is introduced for this purpose. Power control has a dual operation. First, it keeps interference at minimum levels by controlling the power transmitted, keeping in the region of the minimum required for successful reception and thus ensuring an adequate QoS level so that the percentage of dropped calls is kept below the acceptable thresholds. Second, this strategy also minimizes the power consumption at the mobile user (called *User Equipment* [UE] in UMTS) and the base station (called

Node-B in UMTS). Recently, new approaches have been proposed for efficient power control in the case of multicast services for 3G-and-beyond mobile networks [Vlo05]. This study focuses on downlink channels and investigates the performance in terms of transmission power while the number of UEs in a cell and their average distance from the Node B changes (increases).

Implementation of multicast in the downlink direction means that each Node-B needs to transmit at some minimum power for maintaining acceptable *Signal-to-Interference Ratio* (SIR) values for all UEs in the group. Because in a multicast group all UEs receive the same data at any given time, transmission is done to all users in the group simultaneously. This point-to-multipoint transmission uses a common transport channel reaching all UEs in the cell that belong to the specific multicast group. A point-to-multipoint transport channel may or may not have the capability of power control. Point-to-point transmission uses a *Dedicated Transport Channel* (DCH) where power control is operated for each UE. In general, a point-to-multipoint channel requires higher power than does a point-to-point channel. However, as the number of UEs increases, the number of point-to-point channels increases equivalently, whereas for a point-to-multipoint channel, even if the number of users increases, they still use the same point-to-multipoint channel and still consume the same power. Therefore, even though using point-to-point channels may appear to be the most efficient solution for a small number of UEs, the point-to-multipoint solution becomes the best choice for a large number of UEs.

In the UMTS architecture, three transport channels are considered for downlink. DCH is a point-to-point channel with power control. The *Forward Access Channel* (FACH) is a point-to-multipoint channel with disabled power control (i.e., it transmits to all users using a constant power value). The *Downlink Shared Channel* (DSCH) is a point-to-multipoint channel, and it has enabled the inner loop power control, therefore it controls QoS, but has increased overhead in the signaling associated with power control.

Two approaches are considered in the following for the evaluation of the correct scheme for power control in multicast/MBMS schemes: (a) switching between point-to-point and point-to-multipoint channels and (b) using only a point-to-multipoint channel per multicast group in each cell. Each approach has different options in terms of channels usage. These are presented and investigated to show their advantages and shortcomings. In what follows, a CAC scheme has been also proposed for UMTS with MBMS.

2.3.3.1 Switching Between Point-to-Point and Point-to-Multipoint Channels

A basic assumption is that a point-to-point channel (DCH) is used for each UE up to a threshold number of UEs beyond which using point-to-point channels is less efficient than using one point-to-multipoint channel (FACH or DSCH) for the whole cell. The appropriate channel to be used at a specific instant is chosen on the basis of the number of UEs within a cell that belong to the specific multicast group. The point of change is when the sum of powers needed for the

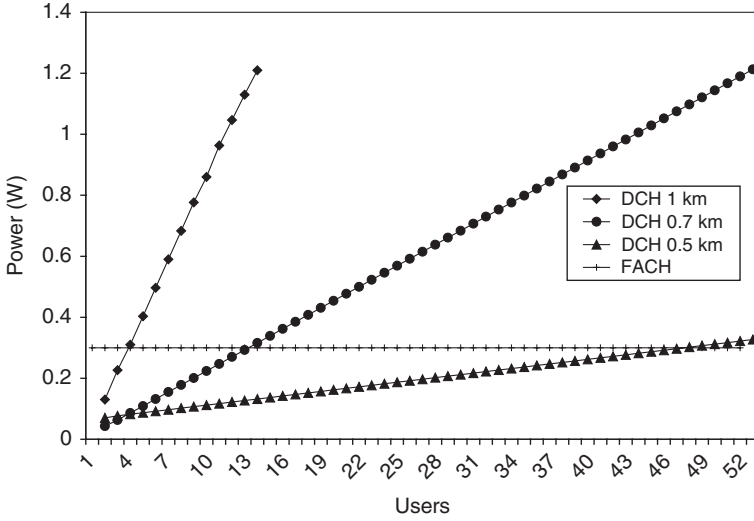


Fig. 2.7 Downlink power for FACH and DCH versus number of users

DCH channels is equal to the power required for the FACH. This number of UEs may serve as a possible threshold for switching between point-to-point and point-to-multipoint channels.

In order to maintain the SIR at acceptable levels, an increase in downlink power with the increase of average distance is expected (see Fig. 2.7). For the FACH, the downlink power is constant, regardless of average user distance or number of users in the cell. This is due to the fact that power control is disabled in FACH. Through simulations (see Fig. 2.7) [Neo06] it has been concluded that in a cell of radius 1 km, the threshold is 47 users if the average distance from the Node-B is 500 m, 13 users if the average distance is 700 m and 3 users if the nodes are at the cell edge (1 km).

2.3.3.2 Using Only Point-to-Multipoint Channels

In this case, only a point-to-multipoint channel is used (FACH or DSCH) per multicast group in each cell. It is expected that as the number of UEs increases, FACH will achieve better performance than will DSCH because power control will add a lot of signaling overhead. Therefore, less overhead and less transmission power should be observed for large numbers of UEs in the case of FACH. However, for a small number of UEs, power control should be enabled to ensure an adequate level of QoS. A research issue is thus to examine if power control for DSCH can be enabled when the number of UEs is below a certain threshold, otherwise it can be disabled.

Simulations prove that the downlink power requirement for DSCH increases with distance (see Fig. 2.8); for FACH of course, downlink power is a constant

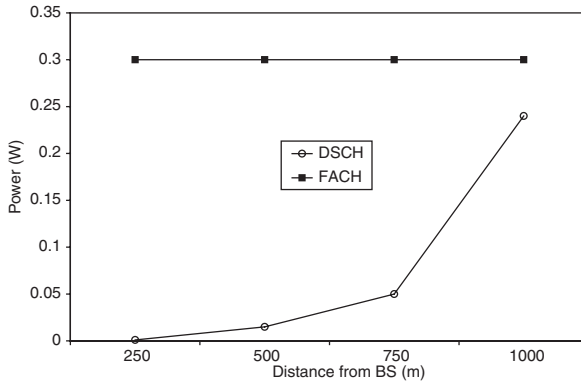


Fig. 2.8 Downlink power for DSCH and FACH for average user distance

irrespective of changing distance. However, up to the distance of 1 km, the performance of DSCH is better than that of FACH in terms of transmitted power, but there is a factor that has not been investigated; that is, the signaling overhead and the uplink power of each UE. There is no such overhead in the case of FACH. The signaling and UE uplink power constitute a large overhead in the case of DSCH, especially when the number of users in the multicast group is large, because a DCH channel is set up for uplink signaling for each user. DSCH may therefore be used for a relatively small number of users. For a large number of users, it may be more efficient to use FACH. However, it is necessary to note that the lack of power control with FACH may cause the calls to have inadequate QoS. Furthermore, it can be observed that DCH and DSCH perform better as the average UE distance from the Node-B is decreased.

2.3.3.3 Call Admission Control in UMTS

A novel hybrid CAC scheme is presented here combining downlink transmission power and aggregate throughput for dedicated or shared connections. The motivation for introducing this hybrid CAC approach is twofold [Ela04, Ela05]: the need to use a representative resource metric on which CAC decisions are based and the necessity to take advantage of the peculiarities of the resource allocation procedure in true multicast environments with connection sharing.

Downlink Power/Throughput-Based CAC (DPTCAC) is the name given to this algorithm, proposed in [Neo06]. It performs admission control based on an estimation of the required downlink transmission power level for the new connection in order to meet its QoS requirements in conjunction with the currently used downlink power levels for ongoing connections and the physical limitation on the maximum transmission power of the base station. This algorithm only considers downlink transmission power (from base station to mobile terminals). Because downlink power levels dominate (i.e., relative to uplink power levels [emitted from mobile terminals]), this simplified consideration is a sufficiently

solid first step toward the overall evaluation of the usefulness of the proposed algorithm. The required transmission power level for a new connection is only an estimate and not an exact value due to the power control mechanism employed in UMTS to regulate transmission power of both base stations and mobile terminals. This mechanism regulates transmission power dynamically according to experienced traffic losses, and therefore under certain circumstances, it is possible for the actual downlink transmission power to surpass or drop below the level estimated during the admission process that is needed to ensure a call's QoS requirements. In addition, user mobility affects the level of required transmission power from the base station in order to maintain the agreed *Signal-to-Interference plus Noise Ratio* (SINR). By considering a worst-case scenario in which the user is located near the cell border, an upper bound can be derived on the estimation of the required transmission power to support the new connection and determine whether this upper limit can be satisfied in the long run. However, the usefulness of this second approximation is tightly connected to the mobility pattern (direction, speed, etc.) of the user. In this study, users are moving fast away from the base station, and the power calculations are based on the instantaneous location of the user. On the other hand, power requirements of low-speed users will not vary dramatically in the short-term, and it is highly unlikely to cause serious fluctuations on base station transmission power.

The estimation of the required base station transmission power level to support the new connection is made using Equation (2.6); the admission decision is then made considering the sum of the estimated required power level and currently used power for ongoing connections against the maximum physical limit of target base station's transmission power (2.7):

$$P_{DL}^* = SINR_t \cdot \left(P_N + \frac{\sum_{j=1}^k P_{DL,j}(r_j)}{N} + \frac{1-\alpha}{N} \cdot P_{DL,0}(r_0) \right) \quad (2.6)$$

$$P_{DL}^* + P_{DL,0} \leq P_{0_max} , \quad (2.7)$$

where:

- P_{DL}^* The estimated required downlink transmission power for the new user (at its current location);
- $SINR_t$ The target SINR that should be met to ensure user QoS requirements;
- P_N The interfering power of background and thermal noise;
- $P_{DL,j}(r_j)$ The total downlink transmission power of base station j (not own base station) perceived at user location (at distance r_j from base station j);
- N Service spreading factor;
- α Own cell downlink orthogonality factor;

$P_{DL,0}(r_0)$	The total transmission power of the target base station perceived at user's location (at distance r_0 from target base station);
k	The number of base stations in the network;
P_{0_max}	The maximum physical transmission power of the target base station.

The hybrid nature of DPTCAC comes into play when admission requests involve multicast services. Power computations described above are performed only once upon establishment of the shared channel for multicast traffic delivery. Subsequent requests for the same service are then only subjected to admission control using shared channel throughput metric in conjunction with shared channel maximum throughput limitation. Concisely, downlink transmission power is taken into account for admission control computations for the initial establishment of a shared channel (FACH or DSCH), while connections running over the shared channel are admitted using the sufficiency of channel's remaining capacity as acceptance criterion.

The performance was compared via simulations of the proposed algorithm against a reference *Throughput-Based CAC* (TCAC) algorithm (analyzed in [Hol01]). Simulation results show a beneficial effect of using DPTCAC on cell capacity without observable degradation of offered QoS.

2.3.4 Interference Coordination in OFDMA-Based Networks⁶

An interesting approach to increase capacity in OFDMA networks is *Interference Coordination* (IFCO) [TD(06)046, Nec07a], where neighboring base stations organize their transmissions to minimize inter-cell interference. IFCO is particularly effective when combined with beam-forming antennas, which additionally allow the exploitation of *Space-Division Multiplexing* (SDM) and thus the transmission to spatially separated terminals on the same frequency/time resource.

IFCO has been an active research area in multihop and mobile *ad hoc* networking, though not so in the area of cellular networks. In [Vil05], the authors study the impact of beam-forming antennas in a multihop wireless network and discuss the implications for MAC protocols. In [Ram89], the authors propose to coordinate broadcast transmissions in a multihop wireless network by means of a greedy graph coloring algorithm, which solves the transmission conflicts of the individual network nodes. In [Jai03], the coordination of transmissions in a wireless *ad hoc* network is considered. IFCO is evaluated by a central entity with full system state information in order to schedule data transmissions of individual nodes at the MAC level. This is done based on a *conflict graph*, which represents critical interference relations between network nodes. The problem of solving the interference conflicts was traced back to the graph coloring problem for example in [Wu05].

⁶ This part was partially published in [Nec07a]. Copyright © 2007 IEEE [Nec07a].

In cellular networks, IFCO has only recently become an active research area, in particular in the course of IEEE 802.16e and 3GPP LTE standardization work (e.g., [R1-051051]). Among the first published studies [Bon05, Liu06], the focus is on a flow-level analysis of the possible capacity gains with intercellular coordination and a static resource assignment policy. Practical approaches discussed in the standardization bodies mostly focus on soft re-use, re-use partitioning, and derived schemes [Sim07]. In [Nec07a], the concept of an *interference graph* is introduced, which represents critical interference relations among the mobile terminals in a cellular network. A simple but efficient heuristic method was used in order to solve the resource assignment problem in combination with the interference graph. Although this approach requires a global device with full system knowledge, it provides important information about key performance parameters and delivers an estimate of the upper performance bound. In [Nec07b], it was extended to an implementable distributed scheme where the base stations communicate with a central coordinator in intervals of the order seconds.

2.3.5 Controlled Contention-Based Access to the Medium in Ad Hoc WLANs

DCF, which is part of the IEEE 802.11 standard [IEEE9], is based on the *Carrier Sense Multiple Access with Collision Avoidance* (CSMA/CA) mechanism. In the DCF scheme, window CW is used by a node in order to control the backoff phase. The backoff time is a random deferral time before transmitting, measured in slot time units. Each node picks randomly a contention slot from the uniform distribution over the interval $[\text{Bound}_{\text{Lower}}, \text{Bound}_{\text{Upper}}]$, where *lower bound* equals 0, and *upper bound* is CW. DCF specifies the minimum and maximum value of CW, $\text{CW}_{\min}$, and $\text{CW}_{\max}$, which are fixed in the DCF standard independently of the environment. Upon each retransmission, CW is doubled ($\text{CW}_{\min} = 32, 64, \dots, \text{CW}_{\max} = 1024$) by the *Binary Exponential Backoff* (BEB) algorithm. Once CW reaches its maximum size ($\text{CW}_{\max}$), it remains at this value until it is reset to the minimum ($\text{CW}_{\min}$). Although CW is doubled, there is always a probability that contending nodes choose the same contention slots, as the lower bound of CW is always *zero*. The adjustment of the upper bound does not consider the network load or channel conditions. This gives rise to unnecessary collisions and packet retransmissions, which lead to energy loss and a shorter network lifetime (if applicable, when nodes are battery-powered). On the other hand, when a transmission is successful or a packet is dropped, CW is reset to the *fixed* minimum size ($\text{CW}_{\min} = 32$). However, a successful transmission and reception of a packet does not say anything about the contention level, but only about picking a convenient slot time from the CW interval. The optimal minimum (maximum) CW value closely depends on the number of nodes actively contending in the network.

To cope with the issue of the *fixed* minimum and maximum CW values, the NEWCAMac [Syl06, TD(06)028] and the NCMac [Rom06] protocols have

been proposed. Both algorithms, *inter alia*, estimate minimum and maximum CW sizes, taking into account the 1-hop active neighborhood. NEWCAMac considers the energy level of the battery for estimating $CW_{\min}$, while NCMac uses this information to adjust not only $CW_{\min}$ but also $CW_{\max}$.

To solve the CW resetting weakness of the DCF mechanism, the *enhanced Dynamic Resetting Algorithm* (eDRA) [Rom07], the *Neighbor and Interference-Aware MAC protocol* (NIAMac) [Blo07], and the combination of both, the *mobile NIAMac mechanism* (mobiNIAMac) [TD(07)028], have been proposed. The NIAMac protocol considers the number of near and farther neighbors and channel conditions, while resetting the CW value. In the eDRA protocol, the CW resetting algorithm takes into account the mobility (fast/slow, increase/decrease) of 1-hop active neighbors during the recovery mechanism and the influence of the number of retransmission attempts. The mobiNIAMac considers the most precious information of both schemes such as mobility of nodes during the recovery mechanism, network load defined by near-far neighbors, and channel conditions. In other words, the NIAMac protocol has been combined with the “mobility” approach of the eDRA protocol.

To decrease the number of collisions caused by the *fixed* values of the lower (always zero) and upper bounds of CW interval, the *selection Bounds* (sB) [Rom07] algorithm was designed in which the backoff timer is randomly selected from the range delimited by *dynamic* lower and upper bounds. The derivation of these bounds considers the number of 1-hop active neighbors and the number of retransmission attempts.

An enhancement of sB, the *Dynamic Energy-Aware Bounds* (dBE) algorithm, seeks to solve the problem of the unequal energy distribution in the network, as with both DCF and sB algorithms some nodes have still a lot of energy when the first node has already died [TD(07)028]. The approach is to use an extra piece of local information, the energy level of battery. The benefit derived from using this extra information in dBE is the reduction of the number of collisions and an improvement of the throughput and energy performance, relative to the sB results.

In a further effort to cope with the contention window and resetting algorithmic issues, dBE and mobiNIAMac algorithms have been combined (i.e., dBE-mobiNIAMac) [TD(07)028]. Figure 2.9 shows a comparison of dBE-mobiNIAMac with DCF 802.11 standard and sB-DRA (i.e., a combination of sB and eDRA algorithms) [Rom07]. In this simulation, the dBE part has been tuned to find the optimal parameters in the trade-off between throughput and energy/lifetime performance. In Fig. 2.9, dBE(γ .5)-mobiNIAMac represents dBE-mobiNIAMac where a suitable value has been used for parameter γ , according to the algorithm shown in [TD(07)028] to determine $\text{Bound}_{\text{Lower}}$ and $\text{Bound}_{\text{Upper}}$ of the modified backoff algorithm.

In Fig. 2.9, dBE-mobiNIAMac achieves significant performance improvements in terms of energy, delay, number of collisions, and throughput, as well as in terms of lifetime, *First Active Node Died* (FND), and *Packet Delivery Fraction* (PDF) measures. Notice that the dBE(3.5)-mobiNIAMac achieves the best

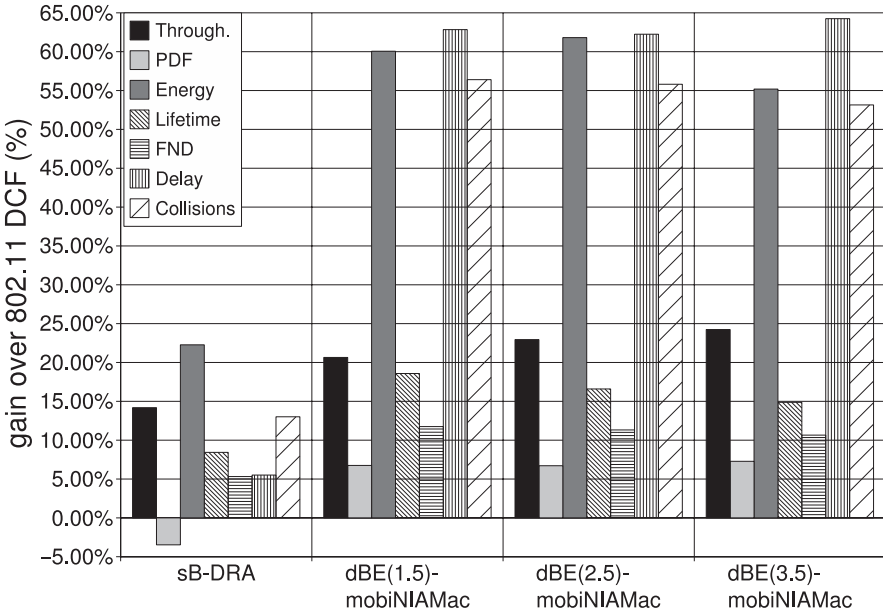


Fig. 2.9 Comparison of dBE-mobiNIAMac with DCF 802.11 standard and sB-DRA

throughput performance, however the worse energy (and lifetime) performance. It seems to be that the most optimal parameter is (2.5), however the choice of a suitable parameter value depends on the QoS requirements.

2.3.6 Multiservice Communications Over TDMA/TDD Wireless LANs⁷

In a TDMA/TDD wireless network, the use of a resource allocation mechanism adapted to the characteristics of connections could increase the performance of the system. Here, a development of a structured set of resource allocation mechanisms in multiservice communications over TDMA/TDD wireless LANs is addressed. One of the first issues to arise is how to make the resource requirements for a particular application available to the Access Point (AP). In the following study, it is shown how network performance improves through the deployment of a set of resource request mechanisms designed to take into account the requirements and characteristics of the specific applications [Del05].

⁷ This work in part derives from the paper published in *Computer Communications Journal*, No. 29, p. 2721–2735, F. M. Delicado, P. Cuenca, L. Orozco-Barbosa, A QoS mechanisms for multimedia communications over TDMA/TDD WLANs, Copyright Elsevier B.V. (2006).

2.3.6.1 Performance Evaluation

Throughout this study, four main traffic types [Del06] have been considered: video [ISO99], voice [ITU92], best-effort [Col99], and background [Kle01]. In order to limit the delay experienced by video and voice applications, the maximum time that a packet of video and voice may remain in the transmission buffer has been set to 100 ms and 10 ms, respectively, on the basis of the standards [Kar00]. A packet is dropped when its delay exceeds its relevant upper bound. In order to evaluate the various resource request mechanisms, a scenario is considered where a third of the *Mobile Terminals* (MTs) is running voice/video applications; another third generates best-effort traffic; and the remaining MTs generate background traffic.

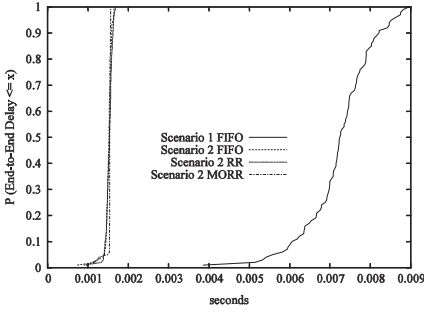
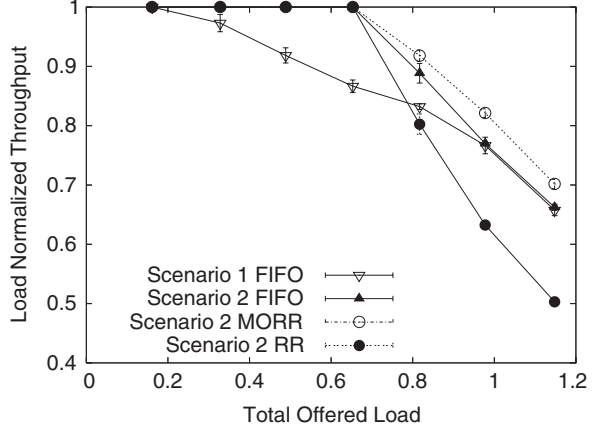
In evaluating the performance and effectiveness of the resource request mechanisms proposed and detailed here, four resource request mechanism types are considered, as detailed below: (i) Type 1, based on contract; (ii) Type 2, based on periodic unicast polling; (iii) Type 3, based on contentions and unicast polling; (iv) Type 4, based only on contentions.

Two sets of simulations corresponding with two scenarios were carried out. In Scenario 1, all applications have to go through a contention-based process when attempting to transmit every resource request packet. In Scenario 2, each application makes use of a different mechanism: voice services use Type 1 mechanism with a guaranteed data rate of 16 kbit/s; video services use Type 2 mechanism with a timer period of 40 ms; best-effort and background traffic use Type 3 and Type 4 mechanisms, respectively. A further objective is to evaluate the performance of bandwidth allocation schemes, such as: *First In First Out* (FIFO), *Round Robin* (RR), and *Minimum Overhead Round Robin* (MORR). The latter is an RR-based scheduler where in each round the scheduler tries to pull all the data of the delivered queue [Del06].

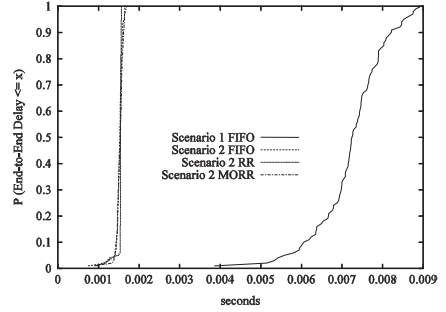
Figure 2.10 represents the normalized throughput (ratio between total input traffic and total granted traffic) as a function of the offered load for both scenarios and all three bandwidth allocation schemes. As seen from the figure, as the load increases, the performance of Scenario 1 badly degrades. This situation can be simply explained as follows. Because the MTs have to go through a contention mechanism to place their requests, as the load increases, the number of collisions in the random access phase will increase dramatically. Furthermore, the fact that the RR bandwidth allocation scheme exhibits the worst results under heavy load conditions is due to the need to dedicate more bandwidth for control purposes. This problem is partially solved by using the MORR scheme at the expenses of penalizing the multiplexing gain.

Figure 2.11 shows the *Cumulative Distribution Function* (CDF) of the end-to-end delay and the jitter for a system operating at full load ($\approx 98\%$). Figures 2.11(a) and (c) show that voice communications are unaffected because the network services guarantee them the required capacity (Scenario 2). In the case of the video traffic, these results show that the MORR mechanism guarantees an end-to-end delay of less than 50 ms to all packets. For the jitter, Fig. 2.11(d) shows that

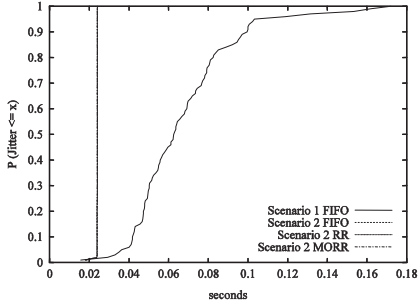
Fig. 2.10 Normalized throughput as a function of the offered load with FIFO, RR, and MORR bandwidth allocation schemes for scenarios 1 and 2 [Del06]



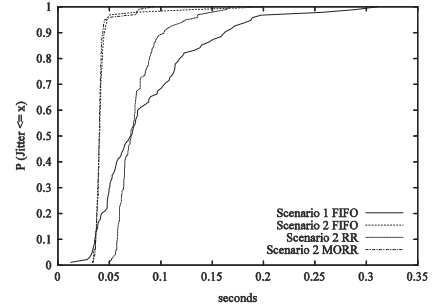
(a) Voice CDF of delay



(b) Video CDF delay



(c) Voice CDF of jitter



(d) Video CDF of jitter

Fig. 2.11 CDF for the end-to-end delay and the jitter for voice and video connections (offered load ≈ 0.98) for scenarios 1 and 2 [Del06]

the 95% percentile of the interarrival times between video frames is 40 ms when MORR or FIFO are used in Scenario 2. This corresponds with the sampling rate of 25 frames/s (i.e., a frame every 40 ms). In other words, 95% of the video frames arrive to their destination in an isochronous way. This is an excellent result that clearly indicates the effectiveness of the proposed mechanism.

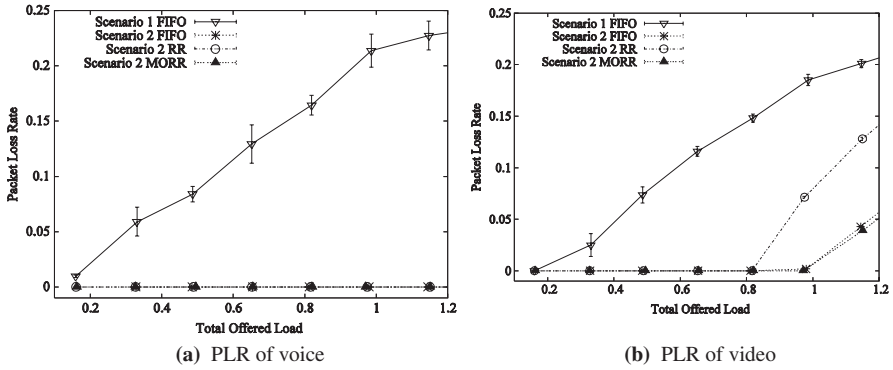


Fig. 2.12 PLR of (a) voice and (b) video connections [Del06]

Figure 2.12 shows the *Packet Loss Rates* (PLRs) for voice and video connections. These losses correspond with the packets dropped as soon as they exceed the maximum allowable queuing delay. In the case of voice connections, Fig. 2.12(a) shows that the losses are completely avoided by statically allocating 16 kbits/s independently of the bandwidth allocation scheme being used. In the case of video connections, the bandwidth allocation scheme plays a major role in their performance. Figure 2.12(b) shows that, for the case when the RR scheme is used, PLR steadily increases when the offered load goes beyond 0.7. Once again, this can be explained by the overhead introduced by this scheme that attempts to multiplex a larger number of connections than the other two bandwidth allocation schemes, namely FIFO and MORR. The use of these two last schemes limits PLR to less than 1% even when the network operates under very heavy load conditions (≈ 1).

In conclusion, the above results show that the use of resource request mechanisms adapted to the requirements of various traffic types is an interesting approach toward the provisioning of QoS guarantees. In evaluating various bandwidth allocation mechanisms, it is shown that it is possible to make use of a simple scheme to reduce the amount of overhead to be introduced into the frame.

2.4 Congestion and Flow Control

Currently-deployed congestion control mechanisms have served the Internet remarkably well as this has evolved from a small-scale network to the largest artificially-deployed system. However, real measurements, simulations, and analysis indicate that the same mechanisms will fail to perform well in the near future as new technology trends (such as wireless access) and services (such as VoIP) modify key characteristics of the network. Moreover, TCP

congestion control is known to exhibit undesirable properties such as low utilization in the presence of large bandwidth-delay products and random packet losses [Lak97, Flo91]. It has also been shown analytically that as bandwidth-delay products increase, TCP becomes oscillatory and prone to instability [Low02].

In this section, some aspects on the performance of currently-deployed congestion control mechanisms are set out and contributions toward the development of new congestion control techniques and protocols, considering the characteristics of the wireless scenario, are presented. In particular, the behavior of Skype voice over IP (VoIP) flows is investigated and, despite the fact that the application implements some sort of congestion control mechanism, it is demonstrated that it still causes problems to the overall network performance. A number of new techniques, architectures, and protocols that overcome the problems of previous approaches, with a focus on solutions that use explicit single-bit and multi-bit feedback, are presented. Quick-Start TCP, a Slow-Start enhancement that uses explicit router feedback to notify the end-users of the allowed sending rate at the beginning of each data transfer, is described and evaluated. A user can then use this information to start sending data immediately with a large congestion window. This is followed by a presentation on the *Open Box Transport Protocol* (OBP), which uses router collaboration to identify the network resources along the path and provides this information to the end-systems. The end-users can then take the necessary congestion control decisions based on the received information. *Adaptive Congestion Control Protocol* (ACP), a dual protocol with learning capabilities where intelligent decisions are taken within the network, is then presented. Each router calculates, at regular time intervals, the desired sending rate that is communicated back to end-users using an explicit multi-bit feedback signaling scheme. The users then gradually adapt the desired sending rate as their actual sending rate. This protocol is shown through simulations to overcome the problems of previous approaches and meet the design objectives. Finally *Fuzzy Explicit Marking* (FEM), a superior active queue management scheme, is presented. It was designed using rule-based fuzzy logic control to supplement the standard TCP-based congestion control.

2.4.1 Assessing the Impact of Skype VoIP Flows on the Stability of the Internet

Internet telephony VoIP applications are seeing quite a phenomenal growth. A good example is Skype, whose explosive growth poses challenges to telecom operators and *Internet Service Providers* (ISPs) both from the point of view of business model and network stability. In this section, using Skype as the VoIP application, the bandwidth adaptability behavior of non-TCP flows (i.e., flows without effective end-to-end congestion control) is examined. The issue is to

understand the impact of these flows, and their growth in the Internet, on the Internet’s legacy services.

In [Flo04], several guidelines for assessing if a flow is harmful for network stability are suggested. One for instance is that a well-behaved flow should not experience persistent drop rates. When this is occurring, it could be indicative of the onset of network congestion. Moreover, TCP fairness is an important issue: persistent drop rates provoked by misbehaving flows would cause bandwidth starvation for TCP flows, a consequence of the TCP congestion control scheme that effectively reacts to packet losses by reducing the input rate. Investigating the behavior of Skype is interesting in this context. This is set out below through the means of different experiments; see [Cic07, TD(07)050] for full experimental details and results. Related work may be found in [Bas06, Che06a]. It is seen there that Skype has bandwidth adaptability, which is a kind of “self-serving” congestion control, but not a “network-serving” congestion control that is without fairness or respect for the needs of best-effort traffic.

2.4.1.1 Investigating Skype Flows Behavior Under Time-Variable Bandwidths

An experiment, which subjects Skype flows to step-like time-varying (i.e., a square-wave) available bandwidth, switching between 16 kbit/s to 160 kbit/s with a 200-s period, is set out. Skype’s adaptation capability to the available bandwidth and also the adaptation transient behavior and duration may be seen in Fig. 2.13(a). It is seen that the sending rate decreases or increases as the link capacity drops or increases, respectively. For example, Skype’s sending rate decreases to just over 16 kbit/s when the available bandwidth drops to 16 kbit/s (e.g., at 300 s) and increases to 90 kbit/s when the available bandwidth increases to 160 kbit/s. Interestingly, the figure shows that the Skype flow takes approximately 40 s to track the available bandwidth, and during this interval it experiences a significant loss rate.

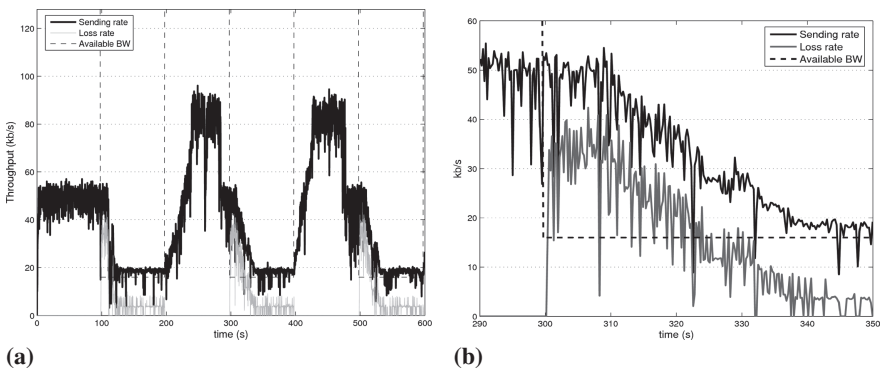


Fig. 2.13 (a) Sending rate and loss rate; (b) zoom around $t = 300$ s

To provide a further insight, Fig. 2.13(b) shows a zoom of Fig. 2.13(a) in the time interval [290 s, 350 s] in order to highlight the effects of a bandwidth drop from 160 kbit/s to 16 kbit/s at $t = 300$ s. The result is that the adaptation kicks in after 10 s, and then the sending rate reduces to less than 20 kbit/s in 30 s. Therefore, it seems that Skype reacts to bandwidth variations with a transient dynamics that lasts 40 s. Loss rates at the higher bandwidth are zero, as this bandwidth is much larger than the load offered. On transition to the lower bandwidth, significant loss rates are initially experienced, interestingly at a rate approximately equal to the difference between the load offered and the 16 kbit/s available bandwidth (i.e., initially around 35 kbit/s). In this experiment, the loss rate eventually settles to around 4 kbit/s. This behavior raises the possibility that Skype may provoke a persistent packet drop rate when a bandwidth reduction occurs.

2.4.1.2 Skype and Fairness in the Presence of Concurrent TCP Flows

A flow is said to be well-behaving if it is able to share fairly the available bandwidth with concurrent TCP flows. In order to address fairness issues, an experiment has been set up where one TCP flow is started over a link with a capacity of 56 kbit/s followed by a Skype flow that starts after 70 s. When the Skype flow is started, its goodput suddenly reaches 40 kbit/s, whereas the TCP flow suddenly decreases almost to zero. In particular, during the period in which the Skype flow is active, the TCP flow experiences frequent timeouts. This result seems to be in contradiction with what has been shown above, where it has been shown that Skype matches the available bandwidth. The reason is that the TCP congestion control reacts to loss events by halving its congestion window, whereas, on the other hand, Skype flows adapt to the available bandwidth slowly. Therefore, even though the TCP congestion control continuously probes for the link bandwidth using its additive increase phase, it is unable to get any significant bandwidth share due to the unresponsive behavior of the Skype flow. Thus, if in the previous Skype-only scenario, persistent losses are experienced in the transition period when the available bandwidth drops, in this scenario it is found that the persistent losses are primarily harmful for concurrent TCP flows; in effect these are at risk of starving out completely.

In the third experiment, four TCP connections are started at different times on a link with constant available capacity in order to see how the Skype flow reacts when a TCP flow joins the bottleneck [Mas99]. This experiment has shown that the Skype sending rate is kept unchanged regardless of number of TCP flows sharing the link, thus confirming the unresponsive behavior of Skype flows in some circumstances.

In conclusion, this experimental investigation reveals that Skype implements a self-serving congestion control algorithm, matching offered load to available bandwidth (this is contrary to what is stated in [Bu06]). However, it seems to lack any fairness attributes for capacity sharing with competing best-effort TCP flows; rather it hogs the available bandwidth, leaving only what it does not need

for other flows. This is definitely not a good sign for legacy services in an Internet environment experiencing exponential growth of Skype-like services.

2.4.2 *Evaluation of Quick-Start TCP*

Quick-Start is a new experimental extension for TCP standardized by the Internet Engineering Task Force (IETF) [Flo07a], which allows speeding up best-effort data transfers. With Quick-Start, TCP hosts can request permission from the routers along a network path to send at a higher rate than that allowed by the default TCP congestion control. This explicit router feedback avoids the time-consuming capacity probing by TCP Slow-Start and is therefore particularly beneficial for underutilized paths with a high bandwidth-delay product, which exist in broadband wide area, mobile, and satellite networks.

In [TD(07)013], a survey of ongoing research efforts on congestion control mechanisms with explicit router feedback is presented. The Quick-Start TCP extension is introduced as one example. Here this extension is detailed and its performance improvement compared with the standard TCP Slow-Start mechanism, both by an analytical model and by simulation results. Initial investigations confirm that Quick-Start can significantly reduce the completion times of mid-sized data transfers. Finally, open issues of congestion control with explicit router feedback are analyzed for Quick-Start.

2.4.2.1 *The Quick-Start TCP Extension*

In Quick-Start, a host can start to send immediately with a large congestion window. Thus, Quick-Start is a performance enhancement for elastic best-effort transport over paths with significant free capacity. Figure 2.14 illustrates a Quick-Start request during TCP connection establishment: in order to indicate its desired sending rate, Host 1 adds a “Quick-Start request” option to the IP header. This option includes a coarse-grained specification of the target rate, encoded in 15 steps ranging from 80 kbit/s to 1.31 Gbit/s. The routers along the path can approve, modify, or reject this rate request. Each router that supports the Quick-Start mechanism performs an admission control and reduces (i.e., reduces the rate) or discards the request if there is not enough bandwidth available.

If the request arrives at the destination Host 2, the granted rate is echoed back, piggybacked as a TCP option (“Quick-Start response”). The originator can then detect whether all routers along the path support Quick-Start and whether all of them have explicitly approved the request. If not, the default congestion control (i.e., TCP Slow-Start) is used to ensure backward compatibility. If the Quick-Start request is successful, the originator can increase its congestion window and start to send with the approved rate, using a rate pacing mechanism (see Fig. 2.14). After one round-trip time, the Quick-Start phase is

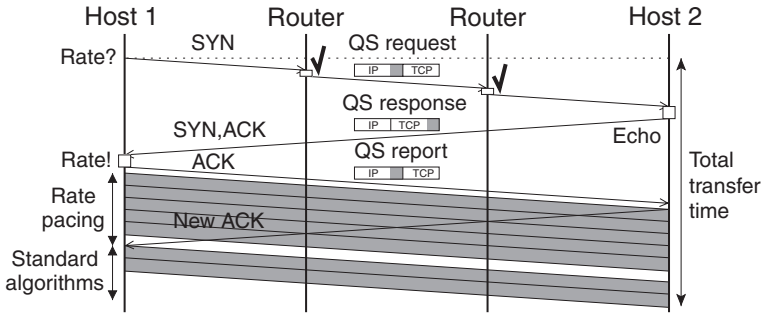


Fig. 2.14 Illustration of a Quick-Start request during the three-way handshake of TCP

completed, and the default TCP congestion control mechanisms are used for the subsequent data transfer.

2.4.2.2 Performance Improvement of Quick-Start TCP⁸

In [TD(07)013], an analytical model has been presented that quantifies the maximum performance benefits of the Quick-Start TCP extension compared with those of standard TCP. This model has also been compared with simulation results. In what follows, a brief summary of these results is provided. A more extensive analysis can be found in [Sch07].

The performance improvement of Quick-Start can be quantified by analyzing the buffer sojourn time T of a given amount of data after connection setup. A couple of analytical models, surveyed in [TD(07)013], incorporate the delaying effect of Slow-Start. With these models, the total transfer time T can be calculated as a function of the available bandwidth of the path, the round-trip time τ , and some further TCP parameters such as the *Maximum Segment Size* (MSS) and the initial congestion window.

In Fig. 2.15, the relative improvement $\eta = T_{\text{Slow-Start}}/T_{\text{Quick-Start}}$ of Quick-Start over standard TCP is depicted for different round-trip times, τ , for 10 MB/s data rate. Both the analytical and simulation results show that Quick-Start can improve transfer times for moderate-sized transfers, in particular if the network latency is high. More specifically, Quick-Start significantly speeds up transfers for a message length in the range of 10 kB to 1 MB. In contrast, Quick-Start is only of limited benefit for very small message sizes, as transfers can be completed in just a few round-trip times anyway. Also, Quick-Start does not significantly improve long bulk data transfers, where Slow-Start is only a transient phase. This study also confirms similar empirical findings in [Sar07].

⁸ Portions reprinted, with permission, from *Proc. 4th IEEE International Conference on Broadband Communications, Networks and Systems (BROADNETS 2007)*, M. Scharf, Performance analysis of the Quick-Start TCP extension. Copyright © 2007 IEEE.

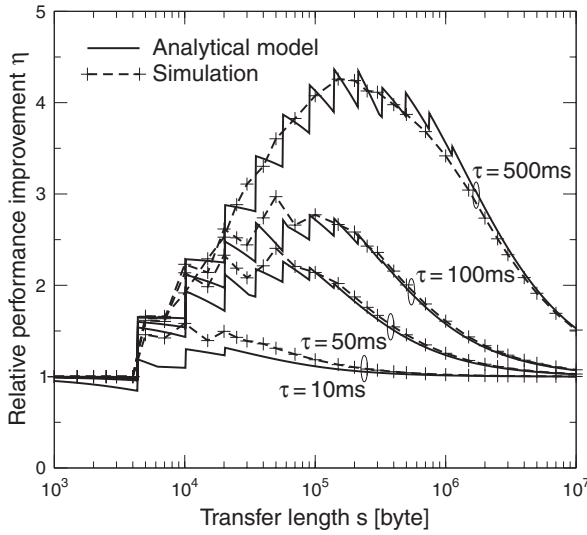


Fig. 2.15 Relative improvement of Quick-Start over Slow-Start (data rate 10 Mbit/s); see Reference [Sch07] (copyright © 2007 IEEE)

In conclusion, the results obtained confirm that Quick-Start can significantly improve transfer times in networks with a high bandwidth-delay product.

2.4.3 Open Box transport Protocol

A new explicit congestion control approach, called *Open Box Transport Protocol* (OBP) is described below. OBP is a cross-layer congestion control protocol, using information and having operations in both the network and transport layers. OBP represents the network path through a small set of variables and continuously provides this information to end-systems. With this information, end-systems make decisions about the use of network capacity and the network congestion state. OBP can quickly adapt to sudden changes in the network, because all transmission rate decisions are supported by the feedback received from the network. Figure 2.16 shows an example of network path with four routers and related links [Lou07].

To represent the network path from one end-system to another one, the following variables are needed by OBP: *narrow link* (link with the most limited capacity), *tight link* (link with the least available bandwidth), *round-trip time*, and *heterogeneous path* (having or not heterogeneous access media along the path, for example wireless links). With the exception of RTT, all the other variables are carried inside fields in the IP header.

In terms of operations, OBP considers that all routers update three variables in each *Acknowledgment* (ACK) packet: narrow link, tight link, and heterogeneous path for all packets forwarded through them.

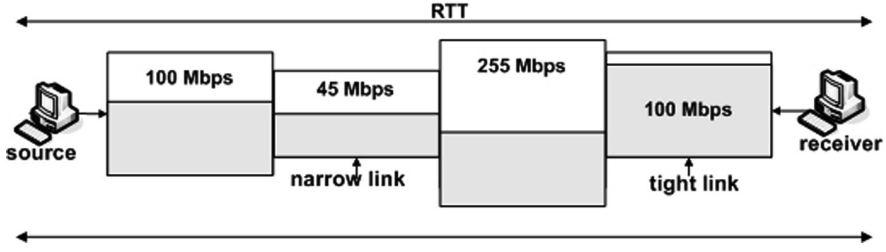


Fig. 2.16 Example of OBP network path representation (copyright © 2007 LNCS-Springer [Lou07])

Unlike other explicit congestion control protocols, the OBP congestion control decisions are made at end-systems, thus freeing the routers for other tasks related to routing and forwarding. The end-systems have the most critical task because they have to take decisions concerning the following elements: performance (whatever the flows size distribution), short flow completion times, fair sharing of available bandwidth, efficient use of high bandwidth-delay product links, capacity to react to sudden changes in the network paths and avoiding congestion.

To address those objectives, OBP uses the following principles: new flows begin with a high transmission rate (this method ensures short completion times for short flows); every time a source receives an ACK packet, the transmission rate is adjusted (these transmission rate adjustments are done to have near-to-zero available bandwidth and, simultaneously, RTT close to the physical minimum, $RTT_{\min}$, i.e., the round-trip propagation delay); the OBP model tries to use efficiently the network path capacity and to avoid congestion (this means that the available bandwidth must always tend to be near zero).

The following equations show how the transmission rate is adjusted in the OBP implementation. First, the initial transmission rate $W(t_0)$ is estimated when the SYN-ACK packet is received on the basis of the available bandwidth $AB(t_0)$, the network capacity $CN(t_0)$ at narrow link:

$$W(t_0) = \alpha \times AB(t_0) + \beta \times CN(t_0) . \quad (2.8)$$

From now on, every time a new ACK packet is received, the feedback information inside the packet is used to make adjustments in the transmission rate. These adjustments are done based on feedback information and based on an equilibrium point. The equilibrium point is updated in multiples of RTT and is computed on the basis of the mean transmission rate during the previous period (this period is equal to an average RTT time).

The transmission rate $W(t)$ depends on the current equilibrium point $EP(k)$, the available bandwidth $AB(t)$, and the network capacity $CN(t)$ at narrow link. Besides that, $W(t)$ is affected by RTT if this value is different from the minimum RTT:

$$W(t) = EP(k) + \frac{EP(k) \times \delta \times AB(t)}{AB(t) + CN(t)} + EP(k) \times \mu \times (RTT_{\min} - RTT) . \quad (2.9)$$

Equation (2.9) allows obtaining transmission rates around the equilibrium point. In other words, if $AB(t)$ is near zero, the $W(t)$ obtained is $EP(k)$. Moreover, if RTT is large, the value obtained for $W(t)$ is less than $EP(k)$. In an extreme case, the $W(t)$ obtained may be near zero, for example if RTT is very high. This behavior protects the network against collapse, because it can instantly reduce the transmission rate to few packets [Flo07b]. Equation (2.9) also enables a quick adaptation to sudden or transient events as it admits changes in the transmission rate whenever an ACK packet is received [Sch07]. Details on Equations (2.8) and (2.9), such as the choice and value of constants used (i.e., α , β , and μ), can be found in [Lou07].

Concerning the equilibrium point, this is updated once per RTT. This way, the equilibrium point is updated with the mean of all transmission rates calculated whenever an ACK packet is received during the previous period:

$$EP(k) = \text{mean} \left[\sum W(t) \right], \text{ during the last RTT.} \quad (2.10)$$

In summary, the formulas used by OBP ensure that the increase in transmission rate is always decided on the basis of the feedback received from the routers. At the same time, the transmission rate can be updated every time an ACK packet is received. Opposite to this behavior, the traditional congestion control algorithms allow the sources to increase the transmission rate without knowing if the network is close to congestion.

In [Lou07], OBP was also positively compared with TCP Reno, *eXplicit congestion Control Protocol* (XCP) [Kat02, Low05], *Rate Control Protocol* (RCP) [Duk05], and TCP Reno with Quick-Start with the request rate equal to 100 kB/s. XCP and RCP were chosen because they also use explicit congestion information to define their transmission rates.

The experimental evaluation has shown that OBP, having the capability to put network state information in the sources, can efficiently use the network bandwidth, keeping the routers' queues near zero occupation. The results have equally shown that OBP can have better performance than that of other congestion control solutions. Moreover, the OBP implementation puts the processing load on the sources side, in opposition to other congestion control approaches, which make congestion control decisions for all flows by the same routers, as is the case for XCP and RCP.

2.4.4 Adaptive Congestion Control Protocol, ACP⁹

Adaptive Congestion Control Protocol (ACP) is a new congestion control protocol with learning capabilities that enable the protocol to adapt to dynamically changing network conditions, to maintain stability, and to achieve good

⁹ This part was partially published in *Computer Networks*, Vol. 51, No. 13 (12 September 2007), pp. 3773-3798, M. Lestas, A. Pitsillides, P. Ioannou, G. Hadjipollas, ACP: A congestion control protocol with learning capability. Copyright © 2007 Elsevier.

performance. The protocol does not require maintenance of per flow states within the network.

The main control architecture of ACP is in the same spirit as that used by the *Available Bit Rate* (ABR) service in ATM networks. Each link calculates at regular time intervals a value that represents the sending rate it desires from all users traversing the link. A packet traversing from source to destination accumulates, in a designated field in the packet header, the minimum of the desired sending rates it encounters along its path. This information is communicated to the user that has generated the packet through an acknowledgment mechanism. The user side algorithm then gradually modifies its congestion window in order to match its sending rate with the value received from the network. The user side algorithm also incorporates a delayed increase policy in the presence of congestion to avoid excessive queue sizes and reduce packet drops.

2.4.4.1 The Packet Header

Similarly to XCP, the ACP packet carries a congestion header that consists of three fields [Les07] as follows. An H_rtt field carries the current RTT estimate of the source that has generated the packet (sender RTT estimate). An $H_feedback$ field carries the sending rate that the network requests from the user application that has generated the packet. This field is initiated with the desired sending rate of the application and is then updated by the ACP protocol at each router and related link the packet encounters along its path. In this way, this field contains the minimum sending rate a packet encounters along its path from source to destination (i.e., desired sending rate requested by the network). The $H_congestion$ bit is a single bit initialized by the user to no congestion (i.e., with a zero value) and set by ACP on a given link if the input data rate at that link is more than 95% of the link capacity (congestion bit). In this way, the router informs its users that it is on the verge of becoming congested so that they can apply a delayed increase policy and avoid excessive queue sizes and packet losses.

2.4.4.2 The ACP Sender

As in TCP, ACP maintains a congestion window, $cwnd$, that represents the number of outstanding packets and an estimate of the current RTT value. In addition to these variables, ACP calculates the minimum of the RTT estimates that have been recorded, $mrtt$. The initial congestion window value is set to 1 packet, and upon packet departure, the $H_feedback$ field in the packet header is initialized with the desired sending rate of the application, and the H_rtt field stores the current RTT estimate. If the source does not have a valid RTT estimate, the H_rtt field is set to zero.

The congestion window is updated every time the sender receives an ACK. When a new ACK is received, the value in the $H_feedback$ field, which

represents the sending rate requested by the network in bytes per second, is read and is used to calculate the desired congestion window as follows:

$$desired_window = \frac{H_feedback \times mrtt}{size}, \quad (2.11)$$

where “size” is the packet size in bytes.

The desired window is the new congestion window requested by the network. $Cwnd$ is not immediately set equal to the desired congestion window, because this abrupt change may lead to bursty traffic; rather, this change is gradually performed by means of a first-order filter. The congestion window is updated according to the following equation:

$$cwnd = \begin{cases} cwnd + \frac{0.1}{cwnd} (desired_window - cwnd), & \text{if } desired_window > cwnd \text{ And } H_congestion = 1 \\ \Pr[cwnd + \frac{1}{cwnd} (desired_window - cwnd)], & \\ otherwise, & \end{cases} \quad (2.12)$$

where the projection operator $\Pr[.]$ is defined as follows:

$$\Pr[x] = \begin{cases} x & \text{if } x > 1 \\ 1 & \text{otherwise} \end{cases}. \quad (2.13)$$

The projection operator guarantees that the congestion window does not become lower than 1.

2.4.4.3 The ACP Router

At each output queue of the router, the objective is to match the input data rate y to the output (link) capacity C and at the same time maintain small queue sizes. To achieve this objective, the router maintains for each link a value that represents the sending rate it desires from all the users' traffic flows traversing the link. For a given flow, the desired sending rate is denoted by p and is updated every control period. Note that y , C , and p are expressed in bit/s. The router implements a per-link control timer. The desired sending rate is updated every time the timer expires. The control period is set equal to the average RTT value, d . Upon packet arrival, the router reads the H_rtt field in the packet header and updates the variables that are used to calculate the average RTT value.

The router measures the input data rate y of each output queue. The router also maintains at each output queue the persistent queue size q that is computed by taking the minimum queue seen by the arriving packets during the last control period. The duration of control period is not constant and is estimated by subtracting the local queuing delay from the average RTT. The local queuing

delay is calculated by dividing the instantaneous queue size with the link capacity. The above variables are used to calculate the desired sending rate p every control period using the following iterative algorithm:

$$p(k+1) = \Pr \left[p(k) + \frac{1}{N(k)} \left[k_i(0.95 \times C - y(k)) - \frac{1}{d(k)} k_q q(k) \right] \right], \quad p(0) = 0 \quad (2.14)$$

where k_i and k_q are design parameters, and N represents an estimate of the number of users utilizing the link, and the projection operator $\Pr[\cdot]$ is modified with respect to (2.13) because it saturates at the C value.

The desired sending rate calculated at each link is used to update the $H_feedback$ field in the packet header. On packet departure, the router compares the desired sending rate with the value stored in the $H_feedback$ field and updates the field with the minimum value. In this way, a packet, traversing from source to destination, accumulates the minimum of the desired sending rates it encounters in its path.

The last function performed by the router at each link is to notify the users traversing the link of the presence of congestion so that they can apply a delayed increase policy. On packet departure, the link checks whether the input data rate is larger than 95% of the link capacity. In this case, it deduces that the link is congested and sets the $H_congestion$ bit in the packet header.

2.4.4.4 Performance Evaluation

Extensive simulations [Les07] indicate that ACP satisfies all the design objectives. The scheme guides the network to a stable equilibrium that is characterized by high network utilization, max-min fairness, small queue sizes, and almost no packet drops. It is scalable with respect to changing delays, bandwidths, and number of users using the network. It also exhibits nice dynamic properties such as smooth responses and fast convergence.

Figure 2.17 shows representative simulation results of the performance of ACP with respect to both TCP and XCP in terms of the achieved goodput in a realistic scenario comprising a relatively few elephant flows (long flows) and a very large number of mice flows (short flows). The network considered contains a single bottleneck link with a bandwidth of 155 Mbit/s and minimum RTT equal to 80 ms. Twenty persistent FTP flows share the single bottleneck link with short Web-like flows. Short flows arrive according to a Poisson process. A number of tests have been conducted changing the mean number of flows entering the network every second in order to emulate different traffic loads. Note that a mean of 500 flows per second is typical in routers experiencing heavy traffic. The transfer file size is derived from a Pareto distribution with an average of 30 packets. The shape parameter of this distribution is set to 1.35.

Figure 2.17 reveals that ACP achieves higher goodput values than do both TCP and XCP. In the case of long flows, the three protocols achieve comparable values, with ACP, however, consistently achieving higher values at all traffic loads. It is also worth noting that TCP achieves relatively low goodput at small

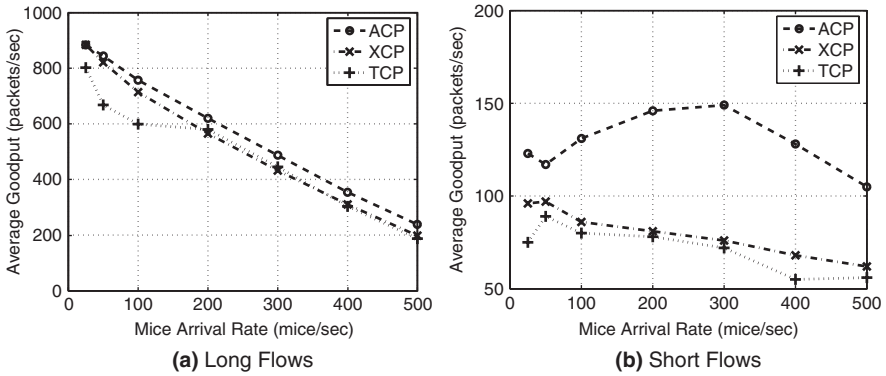


Fig. 2.17 Average goodput for long and short flows for ACP, XCP, and TCP

traffic loads, but as the load increases it achieves almost the same values as XCP. In case of short flows, the superiority of ACP is evident with the goodput exceeding in some cases twice the value achieved by both TCP and XCP.

2.4.5 Fuzzy Explicit Marking, FEM

Active Queue Management (AQM) mechanisms have been introduced at routers to support the standard TCP congestion control, as the wide replacement of the current TCP congestion control approach does not appear to be pragmatic, at this point of time. A number of AQM mechanisms for TCP/IP networks have been introduced in the literature, such as *Random Early Detection* (RED) [Flo93], *Adaptive RED* [Flo01], *Random Exponential Marking* [Ath01], *Proportional-Integral Control* [Hol02] and *Adaptive Virtual Queue* [Kun04]. The interest is toward the ability of effectively controlling the congestion in dynamic, time-varying TCP/IP networks, thus providing QoS, high link utilization, minimal losses, and bounded queue fluctuations and delays.

The proposed fuzzy control methodology for AQM offers significant improvements in controlling congestion in TCP/IP networks under a wide range of operating conditions, without the need for retuning control parameters. In particular, the proposed fuzzy logic approach for congestion control [Chr06], in both best-effort and DiffServ environments, allows the use of linguistic knowledge to capture the dynamics of nonlinear probability marking functions and uses multiple inputs to capture accurately the (dynamic) state of the network. The fuzzy logic-based AQM control methodology better handles the nonlinearity of the TCP network and thus provides an effective control for congestion.

2.4.5.1 Fuzzy Explicit Marking Control System

The proposed nonlinear *Fuzzy Logic-Based Control System* (FLCS), as shown in Fig. 2.18 [Chr06], follows an AQM approach where it implements a drop

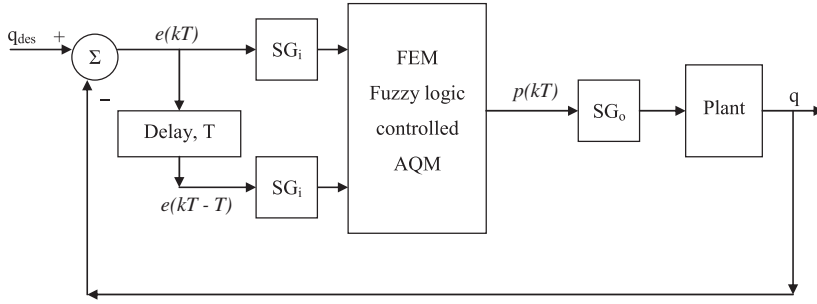


Fig. 2.18 FEM system model [Chr06] (copyright © 2006 IEEE)

probability function and where it supports *Explicit Congestion Notification* (ECN) in order to mark packets instead of dropping them. It uses feedback from the instantaneous queue length and is driven by the error between a given queue reference for the current and previous period. The end-to-end behavior of TCP is retained, with the TCP increase and decrease algorithm responding to ECN marked packets.

All quantities in the system model are considered at the discrete instants kT , with T the sampling period; $e(kT) = q_{des} - q$ the error on the controlled variable queue length, q , at each sampling period; $e(kT - T)$ the error of queue length with a delay T (at the previous sampling period); $p(kT)$ the mark probability; and SG_i and SG_o are scaling gains.

The *Fuzzy Inference Engine* (FIE) uses linguistic rules to calculate the mark probability based on the input from the queues, as set out in Table 2.2.¹⁰ Usually multi-input FIEs can offer better ability to describe linguistically system dynamics. It is expected that in this way it is possible to improve

Table 2.2 FEM linguistic rules: rule base [Chr06] (copyright © 2006 IEEE)

		$Q_{error}(kT - T)$						
		NVB	NB	NS	Z	PS	PB	PVB
$Q_{error}(kT)$	NVB	H	H	H	H	H	H	H
	NB	B	B	B	VB	VB	H	H
	NS	T	VS	S	S	B	VB	VB
	Z	Z	Z	Z	T	VS	S	B
	PS	Z	Z	Z	Z	T	T	VS
	PB	Z	Z	Z	Z	Z	Z	T
	PVB	Z	Z	Z	Z	Z	Z	Z

¹⁰ Table content notations: negative/positive very big (NVB/PVB), negative/positive big (NB/PB), negative/positive small (NS/PS), zero (Z), huge (H), very big (VB), big (B), small (S), very small (VS), tiny (T).

the behavior of the queue by achieving high utilization, low loss, and low delay.

The dynamic way of calculating the mark probability by FIE derives from the fact that according to the error of queue length for two consecutive sample periods, a different set of fuzzy rules (and so inference) applies. The mark probability behavior under the region of equilibrium (i.e., where the error on the queue length is close to zero) is smoothly calculated. On the other hand, the rules are aggressive about increasing the probability of packet marking sharply in the region beyond the equilibrium point. These rules reflect the particular views and experiences of the designer and are easily related to human reasoning processes and gathered experiences. Usually, to define the linguistic values of a fuzzy variable, Gaussian, triangular, or trapezoidal shaped membership functions are used. Because triangular and trapezoidal shaped functions offer more computational simplicity, they can be a good choice for this study.

2.4.5.2 Performance Evaluation

Extensive simulations in the *ns-2* environment [ns207] indicate that FEM satisfies all the design objectives [Chr06]. Specifically, the proposed methodology is able to compensate for varying round-trip delays and number of active flows, as well as in dynamic traffic changes and in the presence of short-lived flows, unresponsive flows, and reverse-path traffic. It shows significant improvement in maintaining performance and robustness with fast system response over a wide range of operating conditions, without the need to (re)tune control parameters, in contrast with other well-known, conventional counterparts such as *Adaptive RED* [Flo01], *Random Exponential Marking* [Ath01], *Proportional-Integral Control* [Hol02], and *Adaptive Virtual Queue* [Kun04].

The performance of the AQM schemes of concern has been investigated under dynamic traffic changes and different traffic loads in a tandem network with multiple types of communication links and multiple congested AQM routers [Chr06]. FEM outperforms the other AQM schemes in terms of better resource utilization and lower delay variation, thus it exhibits a more stable and robust behavior with a bounded delay as shown in Fig 2.19. The other AQM schemes show a poor performance as the traffic load increases, achieving much lower link utilization, and large queuing delays, far beyond the expected value. It is clear that FEM has the lowest variance in queuing delay, resulting in a stable and robust behavior. On the other hand, the other AQM schemes exhibit very large queue fluctuations with large amplitude that inevitably deteriorate delay jitter.

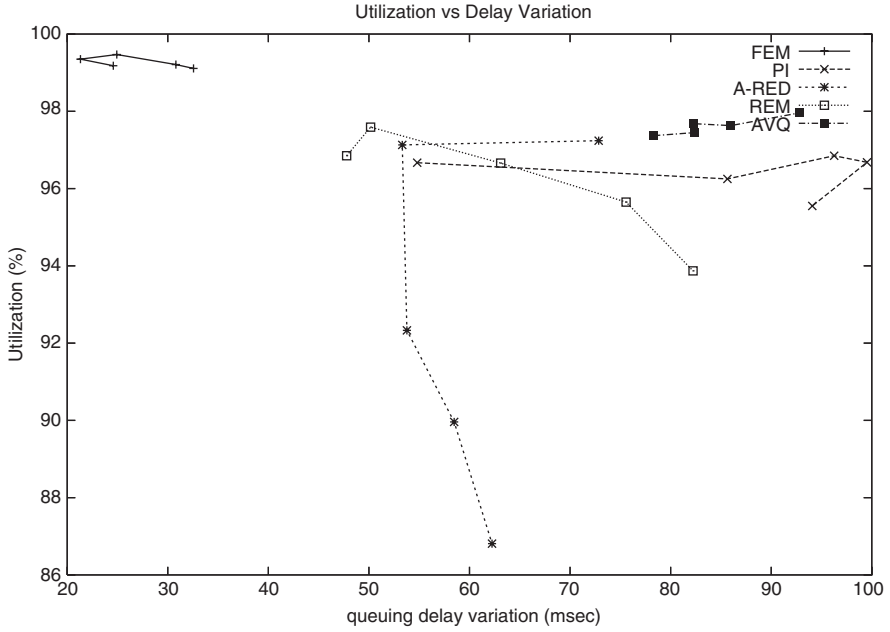


Fig. 2.19 Utilization versus delay variation for different AQM schemes

2.5 Transport Protocols Over Wireless Networks

Understanding the relationships between low-layer performance and the perceived quality at upper protocol layers is of paramount importance to seeking performance improvements for wireless communication systems. Of particular relevance for such a study are terrestrial (e.g., [Alw96]) and satellite wireless scenarios. For the former, the IEEE 802.11 WLAN will be considered here, and for the latter, a *Geostationary Earth Orbit* (GEO) satellite network scenario [Gia06].

This section focuses on the interaction of lower-layer QoS support mechanism with transport layer protocol with the aim to optimize the performance and the efficiency of wireless systems. The interest is here related to the selection of parameters at the MAC layer of IEEE 802.11e to support different traffic flows (including TCP-based applications) and to assess TCP efficiency and the performance of video delivery over asymmetric GEO satellite links.

2.5.1 Wireless Systems

IEEE 802.11 WLANs represent a well assessed solution for providing ubiquitous wireless networking [Man03]. Although nowadays they are widely

deployed, they have two main limitations: (i) the inability to support real-time multimedia applications [Bog07]; (ii) the very high energy consumption due to the wireless network interface card activity [Bog06]. Given that the use of real-time multimedia applications is ever increasing [Wu01], these drawbacks could seriously limit the future development of hotspots based on 802.11. Recently, with the IEEE 802.11e amendment, that standard may now support service differentiation.

Several innovations have been added in IEEE 802.11e [IEEE05]: (i) the HCF; (ii) a CAC algorithm; (iii) specific signaling messages for service request and QoS level negotiation; (iv) four *Access Categories* (ACs) with different priorities to map the behavior of traffic flows with user QoS requirements. The HCF protocol uses a contention-based mechanism and a polling-based one: EDCA and HCCA, respectively. HCCA requires a centralized controller, called *Hybrid Coordinator* (HC), generally located at the AP. The HCF is in charge of assigning *Transmission Opportunities* (TXOPs) to each AC in order to satisfy its QoS needs. TXOP is defined as the time interval during which a station has the right to transmit and is characterized by a starting time and a maximum duration. The contiguous time during which a TXOP is granted to the same station with QoS capabilities (i.e., a *QoS Station*, QSTA) is called service period.

2.5.1.1 Transport Protocol Interaction with EDCA

The EDCA uses distinct traffic classes distinguished in terms of ACs. Each AC has its own transmission queue (in both the AP and the QSTAs) and its own set of channel access parameters. EDCA enhances DCF by introducing a new backoff instance with a separate backoff parameter set for each queue. The scheme of the backoff timer for each AC is similar to the legacy DCF backoff procedure.

Service differentiation among ACs is achieved by setting different values for the following parameters $CW_{\min}$, $CW_{\max}$, AIFS (*Arbitration Inter-Frame Space*), and TXOP limit. If one AC has a smaller AIFS or $CW_{\min}$ or $CW_{\max}$, the corresponding traffic has a better chance of accessing the wireless medium earlier than does traffic from other ACs. Generally, AC3 and AC2 are reserved for real-time applications (e.g., voice or video transmissions), whereas AC1 and AC0 are used for best-effort and background traffic (e.g., file transfer, e-mail). Therefore, the appropriate selection of these parameter values is a challenging task that has to be related to the characteristics of higher-layer protocols, adopted applications, QoS requirements, number of users, and traffic load. Such optimization is one of the aims of this study. It may be noted that according to the standard [IEEE05], by means of the beacon frames, the AP can update QSTAs with new values of AIFS, $CW_{\min}$, $CW_{\max}$, and TXOP limit for the different ACs to cope with varying system conditions. Moreover, $CW_{\min}$ and $CW_{\max}$ must have values belonging to the set $\{2^X - 1$, where X is a number with 4 binary digits}. The value of TXOP limit is a multiple of 32 μ s and varies in

the range $[0, 8160] \mu\text{s}$; a TXOP limit field value of 0 indicates that a single packet is transmitted at any rate for any transmission opportunity.

It is well known that IEEE 802.11e introduces unfairness problems between uplink and downlink flows [Lei05, Cas05]. This is particularly important for the study that will be performed here with VoIP and FTP applications, both characterized by bidirectional flows. In such a case for the VoIP application, a worse voice quality could be perceived by QSTAs, while in the FTP download case, downlink transmission could experience delays with a significant goodput reduction. This means that without a suitable prioritization scheme of downlink flows with respect to uplink ones, bidirectional flows are unbalanced. In this study, a simple approach is proposed where the priority of downlink flows is increased by allocating them an AC with higher priority than that of the corresponding uplink flows; moreover, the contention window size is adjusted according to a simulation approach that is described in what follows.

The system scenario considered here envisages a number of QSTAs, each having a bidirectional VoIP (G.729/A) transmission (*User Datagram Protocol* [UDP]-based traffic with constant bit-rate), together with an FTP downlink (TCP-ACK-clocked) flow from the network. The traffic flows with related directions and ACs mapping are set as follows:

- VoIP in downlink on AC3;
- VoIP in uplink on AC2;
- FTP data in downlink on AC1;
- TCP ACKs in uplink on AC0.

The interest here is to investigate the impact of the different EDCA parameters at MAC layer on the transport layer performance, that is, TCP goodput and VoIP (end-to-end) mean packet delay. The simulation setup uses four ACs, each having AIFS, $CW_{\min}$, $CW_{\max}$, and TXOP values assigned. It is therefore quite complex to investigate joint and individual impact of these parameters on the performance of the transport layer flow and to determine an optimized configuration. In order to restrict the investigation, the AIFS and TXOP values of the four ACs are set to the (default) values defined in the standard. With this, the optimization study then targets the selection of $CW_{\min}$ and $CW_{\max}$ values for AC3. As a first consideration, it may be admitted that these values depend on the number of contending nodes (small windows may increase the number of collisions in the presence of many nodes) and on the degree of prioritization of a traffic flow with respect to another ones (mapped onto different ACs).

In the simulation, using *ns-2*, [ns207], both $CW_{\min}$ and $CW_{\max}$ for AC3 were varied. $CW_{\max}$ was set to allow one expansion of the CW value after a collision. In this study, IEEE 802.11b at 11 Mbit/s at the physical layer is used and a *Frame Error Rate* (FER) of 3% is considered. The results obtained for both TCP and UDP flows are shown in Fig. 2.20. It can be noted that as $CW_{\min}$ of AC3 increases (i.e., $CW_{\min 3}$ in Fig. 2.20), the prioritization of the VoIP downlink flow reduces; correspondingly, the delay for the VoIP uplink flow is lower and the TCP goodput better. It may be noted that TCP does not achieve a good

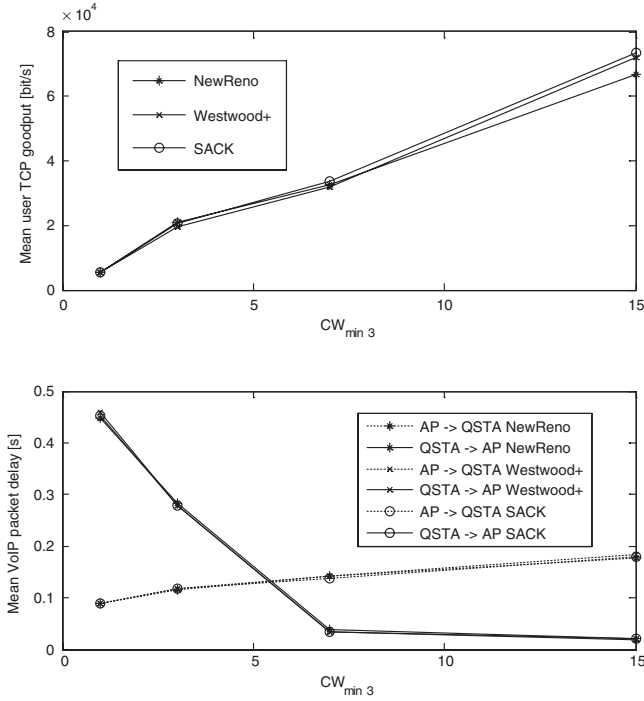


Fig. 2.20 Impact of CW_{min} value for AC3 on the performance of both VoIP and FTP traffic with seven QSTAs and $FER = 3\%$ [Alo07] (copyright © 2007 IEEE)

performance. This is due to the fact that TCP traffic suffers from significant delays arising from low priority, frequent contentions, and packet losses. For all these reasons, TCP flows are unable to widen the congestion window in a proper way to achieve high goodput. Finally, in terms of the VoIP traffic, $CW_{min} = 7$ and $CW_{max} = 15$ may be identified as demarking the best operating range for AC3. No significant performance differences are noted in this configuration for the investigated TCP versions. Other simulations have been carried out to evaluate the impact of the AC1 window selection, and results indicate that increasing CW_{min} for AC1 has no impact on the VoIP traffic that uses higher-priority ACs; whereas there is certain impact on the TCP goodput performance for which a good selection is $CW_{min} = 15$ and $CW_{max} = 1023$ (default value) for AC1.

Similar behaviors of performance parameters to those shown in Fig. 2.20 have been obtained in the presence of three QSTAs, thus confirming the above considerations. Obviously, goodput performance in the case of three QSTAs was much higher due to the reduced time spent in contentions.

In EDCA, MAC-layer transmissions are based on a layer 2 ACK scheme to cope with the uncertainty of collisions and packet errors. The data packet must

be retransmitted if after its transmission a timeout expires before an ACK is received. This mechanism is employed for FTP data packets as well as for VoIP ones. The retransmission scheme means most packet errors due to the channel are hidden from the higher-protocol layers, unless repeated and combined packet errors and collisions are experienced such that the retry limit is reached and the related packet is discarded by the MAC-layer queue. Hence, the retransmission mechanism causes additional delay and the possible drop of some packets, and, in cases of congestion, MAC-layer buffer overflow may occur. These events have a direct impact on the TCP injection rate (i.e., congestion window, *cwnd*, behavior). Hence, at the TCP level some packet drop or buffer overflow events are experienced that need to be recovered with the typical mechanisms (e.g., those used by NewReno, SACK, or Westwood+).

The optimization process studied here could be extended to different situations by means of dynamic AC parameters selection, controlled by AP on the basis of a stored database of optimal parameter settings (and a suitable CAC scheme) for different system conditions.

The results obtained in this study prove that AC mapping issues and appropriate settings of corresponding MAC-layer parameters can have a significant impact on the performance of applications. Moreover, this study shows possible difficulties that may arise in the given IEEE 802.11e QoS-related environment when the both TCP and UDP traffic compete for resources. Interested readers are directed to [STSM(06)002, Alo07] for more details.¹¹

2.5.1.2 Transport Protocol Interaction with HCCA

With HCCA, the HC can start a *Controlled Access Phase* (CAP), during which only QSTAs that are polled and granted with the *QoS CF-Poll frame* can transmit for the assigned TXOPs. The number of CAPs and their starting instants are chosen by the HC in order to satisfy QoS needs of each QSTA. CAP length cannot exceed the value of variable *dot11CAPLimit*, which is advertised by the HC using the beacon frame at the beginning of each super-frame.

It has been shown that HCF needs fine tuning to provide the expected QoS [Bog07]. Besides, using HCF, there is the trade-off between power efficiency and packet delay [Cos05]. To deal with these issues, the interest here is focused on the *Feedback-Based Dynamic Scheduler* (FBDS) proposed in [Bog07] and its power saving extension developed in [TD(05)032, Bog06]. FBDS, which has been designed using classic feedback control theory, exploits HCCA for distributing TXOPs to each real-time flow, by taking into account the queue levels fed back by the QSTA hosting the flow. The WLAN system considered is assumed to be made up of an AP and a set of QSTAs. Each QSTA has N queues, with $N \leq 4$, one for each AC in the IEEE 802.11e standard. Let T_{CA} be the time interval between two

¹¹ This activity has been carried out by Mr. Ivano Alocci (University of Siena) in the framework of a COST 290 short-term scientific mission at the Tampere University of Technology, Finland [STSM(06)002].

successive CAPs. At the beginning of interval T_{CA} (assumed constant), the AP should allocate the bandwidth that will drain each queue during the next CAP. At the beginning of each CAP, the AP is assumed to be aware of all the queue levels q_i , $i = 1, \dots, M$, at the beginning of the previous CAP, where M is the total number of traffic queues in the WLAN system. The latter is a worst-case assumption; in fact, queue levels are fed back using frame headers [IEEE05]; as a consequence, if the i -th queue length has been fed back at the beginning of the previous CAP, then the feedback signal might be delayed up to T_{CA} seconds. The dynamics of the i -th queue can be described by the following discrete-time linear model:

$$q_i(k+1) = q_i(k) + d_i(k)T_{CA} + u_i(k)T_{CA}, \quad (2.15)$$

where, with reference to the k -th CAP, $q_i(k)$ is the i -th queue level at the beginning; $u_i(k) \leq 0$ is the average depletion rate of the i -th queue (i.e., $|u_i(k)|$ is the bandwidth assigned to the i -th queue); $d_i(k) = d_i^s(k) - d_i^{EDCA}(k)$ is the difference between $d_i^s(k) \geq 0$, which is the average input rate at the i -th queue; and $d_i^{EDCA}(k) \geq 0$, which is the amount of data transmitted by the i -th queue using EDCA divided by T_{CA} . The signal $d_i(k)$ is unpredictable because it depends on the behavior of the source that feeds the i -th queue and on the number of packets transmitted using EDCA. Thus, from a control theoretic perspective, $d_i(k)$ can be modeled as a disturbance. Without loss of generality, a piece-wise constant model for the disturbance $d_i(k)$ can be assumed [Bog07]. Because of this assumption, the linearity of the system model described in (2.15), together with the superimposition principle that holds for linear systems, allows the FBDS to be designed by considering a single step disturbance of amplitude d_0 , that is, $d_i(k) = d_0 \times 1(k)$; in fact, a general piece-wise time-waveform can be obtained by superimposition of single disturbance inputs having different amplitudes and starting times.

The extension of FBDS to manage the power-saving proposed in [Bog06], in particular the *Power Save FBDS* (PS FBDS) algorithm [TD(05)032], is discussed in the following. With such an extension, at the beginning of each super-frame, a station using PS FBDS wakes up to receive beacon frames. Then, the QSTA does not transit into the *doze state* until it has received the QoS-Poll frame and the TXOP assignment from the HC. After the station has drained its queue according to the assigned TXOP, it will transit into the *doze state* if and only if its transmission queues are empty. Moreover, a QSTA in the *doze state* wakes up whenever any of its transmission queues is not empty. In this case, after the transition to the *awake state*, the backoff timer for that QSTA is set to zero. As a consequence, the considered QSTA gains the access to the channel with a higher probability than do other stations using classic EDCA.

To test the effectiveness of PS FBDS with respect to FBDS and EDCA, the proposed algorithms have been implemented in the *ns-2* environment [ns207], considering an IEEE 802.11e WLAN scenario with a mix of α MPEG-4 flows, α H.263 VBR flows, $3 \times \alpha$ G.729 flows, and α FTP flows, where α will be referred to as a load factor. Traffic models are the same as those used in [Bog06]. This

load factor has been varied in order to investigate the effect of different traffic conditions on the performance of the considered allocation algorithms. Each wireless node generates a single data flow. According to the IEEE 802.11 standard, in the *ns-2* implementation, T_{CA} is expressed in *Time Units* (TU) equal to $1024 \mu s$ [IEEE05]. In what follows, a T_{CA} of 29 TU is assumed. A data rate equal to 54 Mbit/s for all the wireless stations has been considered. Stations hosting FTP flows do not use any power-saving extension. FTP flows are used to fill in the bandwidth left unused by flows with higher priority.

Figure 2.21 shows the average packet delay experienced by MPEG-4 and G.729 flows as a function of the load factor α . It can be noticed that G.729 flows

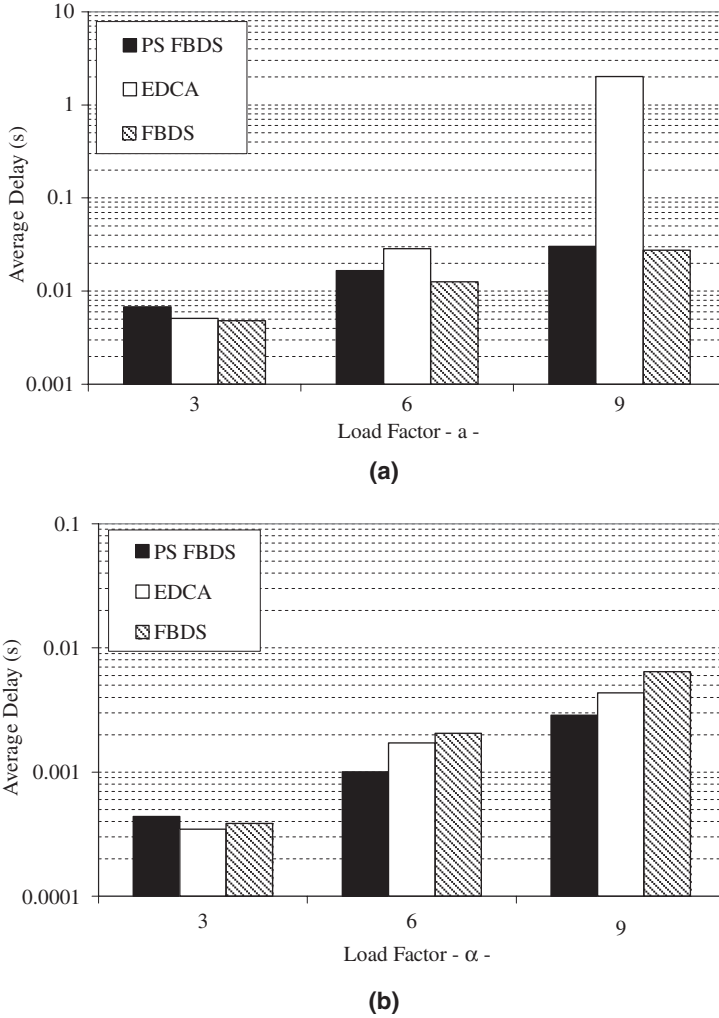


Fig. 2.21 Average delays of (a) MPEG-4 flows; (b) G.729 flows

always get a very small delay. The reason is that voice flows have the highest priority [IEEE05]. Anyway, when FBDS is not used, video flows, having a smaller priority, are penalized. Otherwise, the use of FBDS or PS-FBDS permits obtaining bounded delays regardless of flow priorities. Further simulation results, not shown here, clearly indicate that PS FBDS is able to reduce power consumption with respect to FBDS and EDCA in all scenarios considered; in this study, the RF transceiver IC is the Maxim MAX-2825 802.11 g/a. In conclusion, it has been shown how energy-efficient HCF-based dynamic bandwidth allocation algorithms can be designed for managing real-time services by using a control-theoretic approach, yielding constrained delays for these services without energy efficiency losses.

2.5.2 GEO Satellite Systems

Satellite communication systems are evolving toward the delivery of broadband IP services and are candidates to integrate terrestrial wireless data networks due to their wide coverage and broadcast capabilities. However, satellite networks have limitations, such as long propagation delays and fading channels (e.g., due to meteorological phenomena) that entail higher *Bit Error Rate* (BER of the order 10^{-6} or worse, depending also on transmit energy per bit) values than are normally encountered in terrestrial fixed networks. In order to provide services at a reasonable cost, satellite links exhibit bandwidth asymmetry [Bal02], as they comprise a high-capacity forward space link and a low-bandwidth reverse (space or terrestrial) path.

Media-streaming applications are comparatively intolerant of delay, and of variations in delay and throughput. Furthermore, reliability parameters, such as packet drops and bit errors, usually represent an impairment factor, as they cause a perceptible degradation in media quality. The TCP capability to utilize the satellite link (i.e., efficiency) has not been studied in depth in the case of media delivery. Most related research efforts focus on bulk-data transmission over satellite IP networks and study the corresponding TCP performance [Aky01]. Numerous studies address the limitation of utilizing inadequate resources (e.g., inefficient use of bandwidth) during the Slow-Start phase. Some improvements for TCP in satellite systems include TCP Spoofing, Indirect TCP, increased initial congestion window, Fast Start, and *Selective Acknowledgments* (SACK) [Hen99, TD(06)034].

2.5.2.1 TCP Performance and Media Delivery Over Asymmetric Satellite Links

Here the effects on TCP performance and media delivery due to the presence of satellite links are studied, with a focus on comparing the performance of different congestion control schemes. The bidirectional satellite links are

asymmetric with respect to forward path and reverse path bandwidth, where the forward and reverse paths have the same propagation delay.

If the downlink channel is congested, the sending rate is reduced, as well as in the presence of packet errors on the satellite link. If both upstream links (also identified with uplink or reverse path) and downstream links (also identified with downlink or forward path) are not heavily congested and the sender is able to receive three *Duplicate Acknowledgments* (DACKs) in response to the packet loss in the forward path, *Fast Retransmit* and *Fast Recovery* [Ste97] are triggered. If the sender has not received three DACKs, a timeout event is triggered, followed by an abrupt *cwnd* reduction that diminishes the sending rate and may cause a noticeable interruption in the stream playback. Additionally, the implication where the sender does not receive a number of ACKs, due to a constrained uplink bandwidth or heavy reverse traffic, is considered. According to [TD(06)034], transmission delay variations in the reverse path impact the corresponding transmission periods and degrade the performance of media delivery. Although TCP manages to relinquish the resources allocated when it detects congestion, it is not able to relieve the congestion in the reverse path. Even if the upstream link has deep queues, the reverse channel will become saturated before the downstream link, thus degrading TCP throughput performance in the forward direction. More precisely, the ACKs generated in response to receiving data packets reflect the temporal spacing of these data packets on the way back to the sender, enabling it to transmit new packets that maintain the same spacing. However, the limited upstream capacity and queuing at the upstream bottleneck router alter the inter-ACK spacing on the reverse path, and at the sender. When an ACK arrives at the upstream bottleneck link at a higher rate than the link can support, the spacing is expanded between them when they emerge from the link, enforcing the TCP sender to clock out new data packets at a slower rate. Therefore, the performance of the TCP connection is no longer dependent on the downstream bottleneck link alone; instead, it is throttled by the rate of arriving ACK. As a result, the rate of *cwnd* growth slows down, while certain TCP variants, those that dynamically exploit bandwidth availability by measuring the rate of incoming ACK, may achieve inadequate bandwidth utilization. Hence, reaching uplink capacity poses the highest threat on asymmetric links.

Simulations, using *ns-2*, [ns207], are employed to assess TCP efficiency and the performance of video delivery over asymmetric satellite links. The simulated network configuration is shown in Fig. 2.22. The bidirectional GEO satellite link has 10 Mbit/s downlink and 256 kbit/s uplink channels, with a link BER of 10^{-5} in both directions, supporting N senders (or sources) transmitting MPEG-4 video streams to N receivers (or sinks), together with M FTP senders transmitting to M FTP receivers.

To overcome the standard TCP maximum window size (i.e., 64 kB) limitation, a window scale option is assumed, and the maximum window size is then adjusted to 240 kB. As packet payload size is set to 1000 bytes, a window may accommodate at most around 240 packets. Because the simulated network

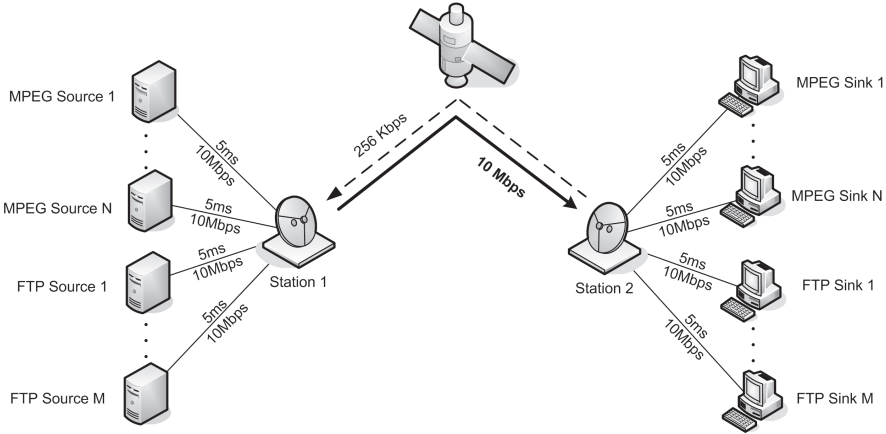


Fig. 2.22 Satellite network simulation topology, with bit rates and propagation delay indicated on links

exhibits an average RTT of 550 ms – of which 540 ms is the fixed full round-trip satellite propagation delay and 5 ms/link is the average terrestrial link delay – the simulation running time was fixed to 200 s, an appropriate time-period for all the protocols.

A wide range of MPEG flows (1 to 50) is simulated, over standard TCP Reno, NewReno, [Flo99], NewReno augmented with SACK, TCP Westwood + (TCPW), [Mas01], and *General AIMD* congestion control (GAIMD; with parameters 0.31, 0.875) [Yan00]. In each case, MPEG traffic shares the satellite link with five FTP connections using TCP Reno. The *goodput* measure is shown in Fig. 2.23(a), whereas the *Delayed Packets Rate*, that is, the proportion of received packets with interarrival times exceeding 75 ms (causing jitter

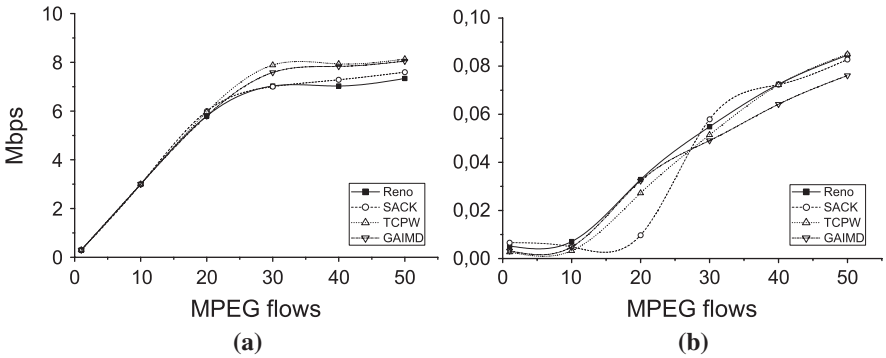


Fig. 2.23 Satellite network link performance: (a) goodput of MPEG flows and (b) delayed packets rate

according to the video streaming delay guidelines), is provided in Fig. 2.23(b). According to Fig. 2.23(a), all TCP protocols are unable to sustain goodput rates close to the bottleneck link rate, despite the relatively large maximum window (i.e., 240 kB). The MPEG connections in each case are affected by link asymmetry, while they are also sensitive to the disturbances caused by competing FTP traffic. More precisely, in the situation of high link multiplexing, the resulting infrequency of ACKs diminishes the sending rate, as *cwnd* is adjusted in response to the incoming rate of ACKs. Furthermore, it was found that Fast Retransmit and Fast Recovery are not triggered, when the upstream link is heavily congested and the TCP sender does not receive three DACKs.

Despite these undesirable implications, GAIMD and, especially, TCPW achieve higher bandwidth utilization, outperforming TCP Reno and TCP NewReno with SACK. Both protocols invoke gentle responses after a packet loss, thus maintaining a higher sending rate. Differently from the initial version of Westwood, TCP Westwood+ computes one sample of available bandwidth every RTT using all data acknowledged in the specific RTT. However, in terms of video delivery, Westwood+ efficiency is not evident, as it delivers a perceptible amount of delayed packets, as shown in Fig. 2.23(b). The protocol responds inappropriately to the variation in the rate of arriving ACKs, as the disturbed inter-ACK spacing reflects the fluctuations in the receiving rate, due to congestion incidents. Unlike TCPW, GAIMD yields satisfactory performance on video delivery for a wider range of flows. The protocol avoids *cwnd* halving by employing a large decrease ratio (0.875), achieving the desired smoothness at the expenses of being less responsive than standard TCP. A comparison between standard TCP Reno and TCP NewReno with SACK reveals that SACK alone is not sufficient to enable high performance as may be deduced from Fig. 2.23(a). However, slight gains are eventually attained, as NewReno prevents coarse timeouts and multiple window reductions, while SACK accelerates the loss recovery phase. Finally, in Fig. 2.23(b), the percentage (or rate) of packets experiencing delays exceeding the streaming video delay requirements may be seen to be not inconsiderable for both Reno and NewReno.

2.6 Cross-Layer Approach

One of the principles for the design of the seven-layer OSI reference model was minimizing interactions between layers and reducing them to the communication through access points between contiguous layers [Zim80]. However, in the past years many works have shown the huge improvement in the performance of wireless systems that can be obtained if certain parameters in one layer are controlled from other layers in the protocol stack [Sri05]. Those techniques have been generally named as *cross-layer* schemes. An extensively accepted example is the presence of adaptive techniques for *Physical* (PHY) and MAC

layers in wireless systems when they can be driven, or partly driven, by performance parameter variations on higher layers.

At present, there is much research on schemes to utilize information from various layers in the protocol stack in order to adapt and optimize the behavior of other layers. In fact, it would be possible to adapt the whole stack in order to globally optimize the system behavior. However, the increased complexity can advise use of direct communication between two layers and locally optimize one of them (usually the lower in the stack) under measures obtained in the other one [Nie06]. In both cases, there is the need of certain information exchange between layers that can follow the same path than data through a service access point or can be done through an external entity. Both schemes for structuring cross-layer interactions are set out in Section 2.6.1.

Providing seamless end-to-end QoS for packet data services over integrated wired and wireless networks represents a significant challenge. Moreover, the mix of heterogeneous physical layer environments, which exhibit different latencies and data rates, leads to a complex scenario from a performance evaluation and optimization standpoint. Two distinct approaches toward QoS are addressed here. The first adapts the streamed content to the current network conditions at the end terminals and is called *end-to-end QoS control* [TD(07)014]. The second offers network support for video streaming services and is called *network-centric control* [TD(07)012]. They are dealt with here in Sections 2.6.2 and 2.6.3, respectively.

Local adaptation in one layer due to changes in the QoS of other layers is being introduced in most current wireless designs such as HSDPA, IEEE 802.11e, and so forth. A number of proposals for current technologies, involving different layers of the protocol stack (PHY-App, PHY-MAC, Link-App, etc.), are detailed in condensed form in Section 2.6.4.

2.6.1 Cross-Layer Performance Control of Wireless Channels

As previously stated, in order to allow cross-layer interactions, certain information exchange between layers is necessary. Cross-layer signaling implementation approaches may be categorized as in-band or out-band. For *in-band cross-layer signaling* [Sud01], two layers directly exchange information used by the performance control entities implemented at each layer through which layer parameters may be dynamically controlled. In *out-of-band cross-layer signaling* [Che02], layers operate on data as in the OSI layered architecture, but they export their current operational parameters to a certain external performance control entity via a predefined set of interfaces. This entity is able to *globally optimize* performance and then distributes information for adaptation of controllable parameters at different protocol layers. On the other hand, in-band signaling only allows *local optimization* of layer parameters.

Global optimization can result too complex or unfeasible, depending on the number of parameters to be optimized. A suboptimal hierarchical optimization would be more manageable. In that sense, the protocol stack can be logically divided into three groups of protocols. The first group consists of an application itself. The second group is formed by transport and network layers. Although in general there is correspondence between the application traffic class and the protocols at those layers, certain usage parameters can be adapted, depending on the application operation. The third group is formed by data-link and PHY layers, usually specific for a given technology.

The proposed structure for a cross-layer performance control system for a wireless access network following the out-of-band model is shown in Fig. 2.24 [TD(07)043]. For all session instances, the *Cross-Layer Performance Optimization Subsystem* (CPOS) monitors the current state of the application, the current state of the wireless channel, and protocols parameters at data-link and physical layers to determine performance parameters such as frame loss rate, frame delay, delay variation, and so forth, and decides which actions on which protocol parameters on which layers should be taken to provide the best possible performance for a given application at the current time instant (e.g., rate of the codec for video and audio applications, the buffer space at the data-link layer, and PDU size at different layers). CPOS is composed of three major components: the *real-time Channel Estimation Module* (rt-CEM), the *real-time Traffic Estimation Module* (rt-TEM), and the *Performance Evaluation and Optimization Module* (PEOM). The rt-CEM is responsible for modeling the

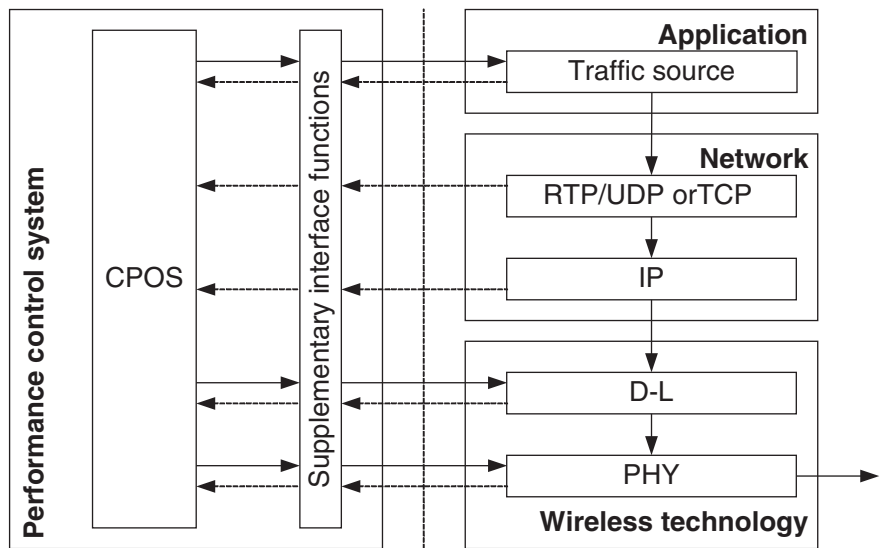


Fig. 2.24 A generic implementation schematic for an out-of-band cross-layer performance control system for a wireless-based network

wireless channel using measurements from data-link or PHY layers such as SNR, bit error rate, or frame error rate. The rt-TEM performs the same functions for traffic observations from application layer.

The cross-layer operation of the whole system is as follows. When a change is detected, the current wireless channel and traffic models are parameterized in the respective modeling blocks and then fed to the input of the optimization PEOM module; this module decides whether performance can be improved under the new conditions and, if so, the new parameter values of the protocol are computed and fed back to the layers themselves.

2.6.2 Cross-Layer Over Wired-Wireless Networks for End-to-End QoS

In this section, an end-to-end QoS model is used to evaluate the performance of data services over wireless-wired networks considering the cumulative performance degradation along each network element and protocol layer [TD(07)014, Gom07].

The radio interface in the wireless domain consists of a generic variable-rate multiuser and multichannel subsystem that is an abstraction from the details of the physical multiplexing technique,¹² considering only parameters that represent the channel time correlation (i.e., variation of signal level in time) and the correlation between channels (i.e., dependence between consecutive channels such as subcarriers in OFDM or slots in TDM) [Ent07]. For Adaptive Modulation and Coding, tracking of the variable quality of each channel is enabled. The model studied here also includes *Robust Header Compression* (ROHC) and retransmissions.

The wired domain has been modeled with three consecutive nodes that are interconnected by links having corresponding capacities over-dimensioned compared with radio link capacity. Each node includes one buffer of 32 kB per user. UDP, TCP, and *TCP-Friendly Rate Control* (TFRC) protocols have been studied at the transport layer. For UDP, the application throughput¹³ can be derived from the estimation on the degradation in the lower layers. For TCP, the application throughput behavior depends on the TCP implementation: TCP Reno has been evaluated following the approach in [Pad98] as a function of the loss rate, the round-trip time, the maximum window size (selected as 16 kB), and the number of packets acknowledged by an ACK. TFRC has been modeled as described in [Flo03].

The UDP-based streaming throughput achieved at different layers is presented in Fig. 2.25(a) as a function of the cell traffic load (i.e., total

¹² For example, by *Time Division Multiplexing* (TDM), *Orthogonal Frequency Division Multiplexing* (OFDM), *Code Division Multiplexing* (CDM), or *Space Division Multiplexing* (SDM).

¹³ Correctly received information bits per seconds at application layer.

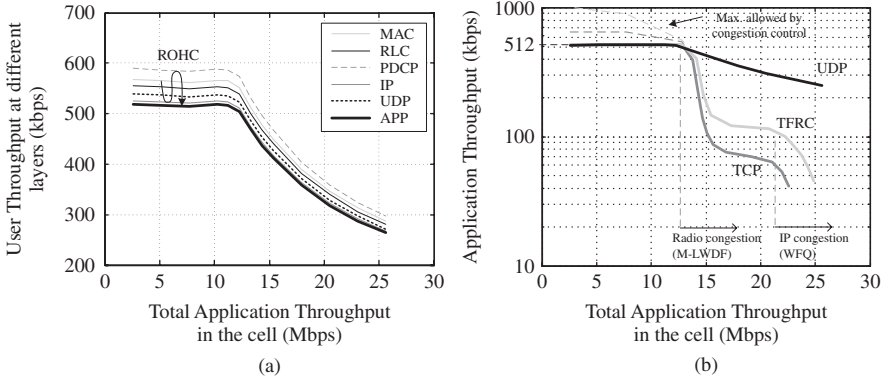


Fig. 2.25 (a) User throughput at different layers for UDP-based streaming with ROHC and (b) throughput comparison at application layer for different transport protocols

throughput generated by all users in the cell). *Modified Largest Weighted Delay First* (M-LWDF) has been selected as the multiplexing algorithm at the radio interface because it considers both the instantaneous channel quality and the queuing delay in the user's priority computation. Application source rate per user is equal to 512 kbit/s. It can be seen that mean MAC layer throughput performance is rapidly degraded (and the other layers in lockstep with it) above a certain critical traffic load (about 12 Mbit/s), because the radio multiplexer is unable to allocate the required resources to all the users. Because of ROHC, *Packet Data Convergence Protocol* (PDCP) layer may achieve a higher throughput than do lower layers. Throughput at upper layers only suffers from RTP/UDP/IP header overheads.

Worse results seem to be obtained for TCP and TFRC (see Fig. 2.25b). However, a particular total throughput corresponds with a different number of users for different transport protocols. Above the critical load point, UDP continues sending at the average codec rate, producing a high loss rate due to overflow in the queues. On the other hand, both TCP and TFRC decrease their sending rates as queue occupancy grows and application data is temporarily stored at the streaming server, and the application may carry out corrective actions. This intrinsic advantage of protocols with congestion control mechanisms (i.e., TCP and TFRC) versus UDP may be a decisive factor in order to select the transport protocol.

The end-to-end QoS model here presented has been valuable for providing performance estimations along the protocol stack, which includes cross-layer interactions to enhance the overall QoS. Although cross-layer design is primarily applied to radio layers (e.g., link adaptation and multiplexing), the proposed model has permitted investigation of interactions with higher-layer protocols.

2.6.3 Network-Centric Methods to Improve Video Streaming QoS

A challenging issue is the provision of video streaming services with a QoS sufficient to meet customer satisfaction. The wireless contribution to quality degradation (e.g., on the final downloading access link) has the potential to be the significant defining part. It is therefore desirable to take suitable measures, additional to those on network connections that do not include wireless links, especially on the access network link. In this section, several network-centric methods are described referring to current technologies [TD(07)012].

Resource reservation (IntServ approach) provides requested QoS by means of the end-to-end resource assignment to a certain traffic flow for the whole session duration. In wired networks with wireless links, both the current capacity and the utilization of wireless links are time-varying values, so that it is difficult to provide absolute end-to-end QoS guarantees.

In the case of *traffic prioritization* (DiffServ approach), the network traffic is classified and different traffic classes are treated unequally in the network elements. The QoS guarantee is relative and realized hop-by-hop. QoS-aware scheduling algorithms include static prioritization and the “early due date” EDD scheme.

In wireless networks with variable rate (due to adaptive modulation or retransmissions), assignment of equal amount of time to different users would result in an unequal treatment. With a channel-aware- (i.e., opportunistic) and class-aware-based scheduler, it is possible to use in the best way the channel capacity [TD(05)045] (i.e., *optimized resource usage*). Channel-awareness means that the scheduler tries to increase channel utilization with simultaneous interuser fairness in a multiuser scenario. Class-awareness means that the scheduler should be aware of traffic class QoS requirements.

Another possibility to reduce or even to avoid transmission errors is video characteristics *adaptation* to time-varying transmission resources (bandwidth) [Fel06]. One approach for this, dynamic *Rate Shaping* (RS) [Cha05], can be visualized as a filter (shaper) that produces an output video stream by changing the input stream according to the current data rate constraints. Another approach, based on RS methods, is video-aware data dropping at network layer, which can be classified as an AQM mechanism [Orl07]: less important or “too late” data are dropped. The objective of *Rate Control* (RC) [Ele95] is to permit the content encoder/server to change the video data rate according to the available resources (bandwidth). In this cross-layer action, video application must be able to monitor/get and process the information about current network conditions [Hem99]. Although the video adaptation is done by the server host, appropriate network nodes are necessary in order to support network monitoring and adaptation decisions.

2.6.4 Application of Cross-Layer Cooperation to Current Technologies

2.6.4.1 VoIP Over Multirate WLANs

In multirate IEEE 802.11-based WLANs, sporadic rate changes, due to the use of link-adaptation mechanisms, can occur in the transmission between a *Mobile Node* (MN) and the AP. Although such rate changes only affect directly the wireless link between one MN and the AP, they impact the quality of all the other active calls [Heu03] because any reduction (increase) in the wireless link rate is equivalent to a reduction (increase) in the whole available bandwidth. In such a scenario, providing the required QoS to VoIP calls could be achieved by using a joint CAC mechanism (MAC-layer) and a VoIP codec selection algorithm. As the latter is an application layer protocol action, the harmonization of both these actions requires a cross-layer approach.

Admission control ensures the stability of the IEEE 802.11 system when new calls/flows arrive and tries to maximize the channel utilization and to guarantee the requirements of all accepted traffic flows. One of the challenges of the CAC scheme is to predict the future system state using current system information. The estimation can be performed based on predesigned mathematical models [TD(06)012] or based on current measurements. In the first case, the models have to be specific for IEEE 802.11 and usually are parameterized simply by the traffic profiles. However, their main problem is that they require solving complex nonlinear computational models to be able to predict accurately the future system state. On the other hand, measured-based prediction is reactive as the new flow has to be already active to decide if it can affect negatively the other active flows.

Additionally, MAC parameters could be modified dynamically in order that the transmission resources would be shared properly and adaptively among all active flows [TD(07)015]. Usually, the algorithm to tune MAC parameters can be integrated into the CAC scheme, as parameters are normally changed when a new flow arrives/departs or when a rate change is detected.

A goal of the VoIP codec selection algorithm is to try to maintain the bandwidth consumption of each call approximately constant despite wireless rate-changes by adjusting the VoIP codec to the channel conditions. The idea is equivalent to the well-known *GSM Adaptive Multi-Rate*, a multirate codec that enables codec rate reduction in reaction to deterioration in channel conditions so as to maintain good or acceptable speech quality for voice calls under different channel conditions [Lun05]. An example of a set of policies based on the *constant relative bandwidth consumption* design criterion (to mitigate the multirate effect) is shown in [Bel06]. In a multirate shared channel, however, it is the behavior of other nodes that is causing the apparent deterioration in channel conditions and hence having an impact on the QoS performance being experienced by the others. This aspect is considered in codec rate adjustment and variation of packetization interval to lower the quality of some of the

existing calls in the cell in order to allow new ones (with main focus on handover calls – handover procedures will be dealt with in Chapter 4) to enter [Che06b]. Hence cross-layer cooperation may be used to engineer solutions for better utilization of wireless access networks.

In the following example, adopting cross-layer cooperation to achieve and maintain acceptable QoS conditions in mixed VoIP and elastic (i.e., with no constraint on delay) traffic scenarios, two modules are necessary: (i) a *Codec Adaptation Algorithm* (CAA), which detects voice quality deterioration from real-time information gathered from the system and proposes a new codec algorithm or codec rate, more efficient for the new cell conditions (see [TD(07)018] for details); and (ii) a CAC scheme, which decides whether to accept or reject new VoIP calls and data flows based on the current system state, the information carried in the admission request transmitted by the MN, and the information provided by the VoIP codec selection algorithm. Moreover, the CAC scheme is able to set the parameters in order to increase the protection and the required QoS for real-time flows, while it tries to minimize the throughput reduction for best-effort flows. The information needed by CAA is gathered from the RTCP packets, providing basic QoS metrics, such as delay, jitter, and packet loss, as well as from the MAC layer, which is enhanced to inform about rate changes when they happen.

In order to evaluate the combined solution for IEEE 802.11e, a flow-level simulator has been used [TD(07)018] that implements the above-proposed mechanisms (see Section 2.5.1 for details on 802.11e EDCA).

The basic benefit of using adaptive policies (i.e., CAA + CAC) may be seen in Fig. 2.26, where the *Grade of Service* (GoS) value (i.e., a sort of blocking probability computed as the sum of the average call blocking probability with weight 10% and the call dropping probability with weight 90%) is plotted. The lower GoS when using the adaptive solution compared with the GoS using only the G.711 and the G.729 codec may be observed. This is caused by the reduction of the dropping probability on rate changes. Note that in this example, the codec adaptation is not used to allocate space for new calls.

Hence, as the results of this experiment indicate, the combination of CAA and CAC yields a better utilization of network resources for a mixed VoIP and elastic traffic scenario than does either on its own, while preserving an acceptable call quality throughout the duration of the call [TD(07)042].

2.6.4.2 H.264 Video Streaming Through DiffServ IP Networks to WLANs

In this section, a practical implementation of an end-to-end H.264 video streaming solution over a wired-cum-wireless QoS-enabled network architecture (combining a WiFi and a fixed network segment) is proposed using a cross-layer architecture based on application, network, and MAC layers [TD(07)021]. A new DiffServ *Per Hop Behavior* (PHB) is presented, suitable for real-time traffic packets having different drop precedence values.

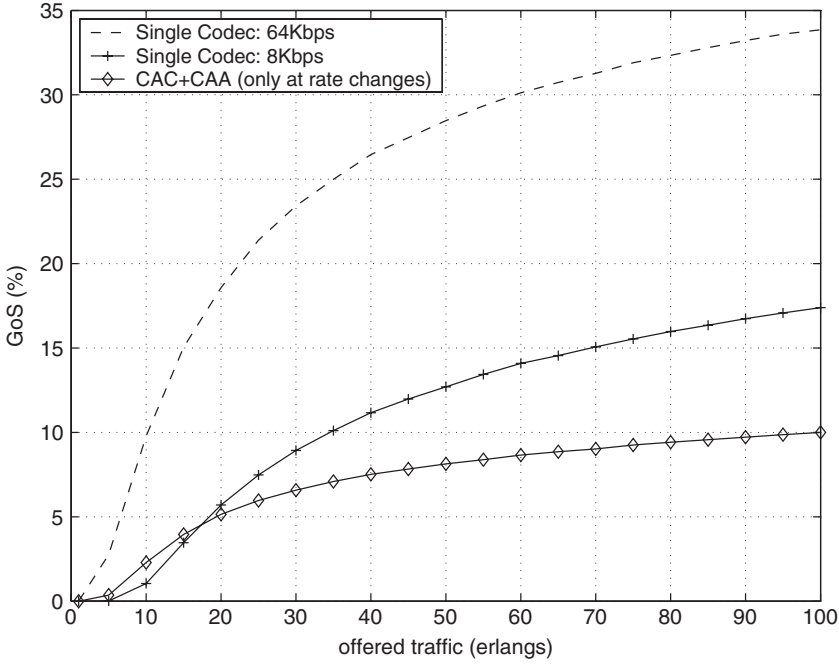


Fig. 2.26 Flow-level performance results for the joint admission control and VoIP codec selection algorithm

H.264 specification includes a *Network Abstraction Layer* (NAL) responsible for the encapsulation of video data into entities suitable for a variety of transport layers or storage media. An *NAL Unit* (NALU) consists of a one-byte header followed by a bit string that contains fixed size picture parts, called *Macro Blocks* (MB). The *Nal_Ref_Idc* (NRI) field in the NALU header specifies the priority of the payload. Video MBs are grouped into partitions with decreasing order of importance: A (headers), B (intrapartition), and C (interpartition). In addition, a slice representing the *Instantaneous Decoding Refresh* (IDR) pictures is generated and the *Parameter Set Concept* (PSC) carries the most important information, relevant to more than one slice [ITU05].

In [Kse06], the authors propose a cross-layer architecture for robust video transmissions over IEEE 802.11e using H.264 (see Section 2.5.1 for details on 802.11e EDCA). In their solution, through the NRI field value, each NALU containing bits from a specific partition (PSC, IDR, A, B, or C) is mapped into an AC of IEEE 802.11e in the range AC1 to AC3 that corresponds with its importance (e.g., PSC is mapped to AC3, highest priority class).

However, including the network layer into the architecture is necessary because in any network containing IP nodes, packets traveling through subsequent routers do not preserve their priority information. A cross-layer architecture solution is possible that extends the work mentioned above by

additionally taking into consideration the network layer in order to propagate the video-related QoS information to the whole network. Thus, besides the AC at the MAC layer, it is possible also to map NRI information extracted from the NALU header to the *Diffserv Code Point* (DSCP) field values at the IP network layer. Mapping of partition classes (A, B, C, IDR, or PSC) to DSCP in current DiffServ classes presents some shortcomings. Assigning Expedited Forwarding (EF) behavior to video streams can cause starvation of other flow aggregates. Moreover, excessive EF traffic in the core network will produce large packet drops with no protection against elimination of important packets [Dav02]. On the other hand, drop priorities for Assured Forwarding (AF) PHBs are usually implemented with a form of RED that can lead to discarding important multimedia packets instead of less important ones [Hei99]. In order to overcome these limitations, a different PHB for multimedia traffic with drop priorities, called *Multimedia* (MM) PHB, is proposed. The MM PHB is similar to EF but additionally employs a strict drop precedence scheme. In this way, important packets have better chances to survive the end-to-end journey.

Experiments to test this hypothesis were designed as set out in the following. A DiffServ *Edge Router* (ER) was configured using Linux QoS mechanisms to police traffic according to the DSCP value of each packet. Traffic conditioning was implemented with four policy filters combined with a DSMARK queue discipline that just marks packets using the DS field. Out-of-profile traffic from a class is re-marked and sent to the lower-priority class. The fourth filter, associated with best-effort traffic, is used to discard out-of-profile packets. The cross-layer architecture was implemented using a combination of open source software [TD(07)021, TD(07)037]. At the application layer, *VideoLan Client* (VLC) open software [Vlc07] has been modified to analyze each NALU and sets the socket's SO_PRIORITY value for the current packet according to the NRI field; at the network layer in the source node, a DSMARK queuing discipline simply translates the SO_PRIORITY value (and indirectly the NRI value) to DSCP without shaping or policing traffic; at the data link layer, the MadWifi WLAN Linux driver for Atheros chipsets was modified to implement DSCP-to-AC mapping.

A H.264-encoded Foreman sequence was sent to the destination through the ingress router. The total rate was limited to 1.1 Mbit/s to enforce reclassification and dropping. The edge router discarded excess packets according to their DSCP set by the video source. A second experiment was performed with similar setup, but, instead of priority dropping, packets were discarded randomly. Image quality was compared using *Peak SNR* (PSNR) for each video frame, and *Average PSNR* (APSNR) was computed for both video sequences received at the destination node, relative to the source Foreman sequence. The APSNR value obtained for the first experiment was 48.72 dB, that is, more than 17 dB higher than that obtained in the second (random drop) case (which was 31.13 dB). The experiments demonstrate the beneficial effects of DSCP-based policing relying on cross-layer information.

2.6.4.3 UMTS/HSDPA Queue Management for Video Transmission

In this section, queue management mechanisms at frame and packet level are considered for a UMTS *Radio Access Network* (UTRAN) with HSDPA extension taking into account the deadline constraints of the video service [TD(07)012].

An MPEG-4-encoded video consists of a number of interdependent frames. If a video streaming session uses an IP-based network as transport infrastructure, the video frames are sent by means of IP packets whose maximum size is often smaller than the average size of a video frame. For this reason, two different IP queue management approaches can be identified: the *packet-based* one and the *frame-based* approach. Whereas the packet-based scheme considers every IP packet as an independent unit, the frame-based approach considers a video frame as indivisible unit during buffer management decisions and actions.

There are three types of frames in an MPEG-encoded video. If an I- or a P-frame is lost, other video frames (P or B) depending on the lost frame cannot be decoded properly. On the other hand, there are no dependencies on B-frames. Consequently, buffer management schemes can be with or without data differentiation. Buffer management decisions with data differentiation depend on the received data (i.e., frame types and their priorities).

Most of today's network elements apply a simple *packet-level drop-tail* FIFO buffer management strategy, where newly arriving IP packets are dropped if the queue is full. The *drop-head* strategy [Orl07] drops those data units that reside longest in the queue (i.e., suppressing data that may arrive too late at the client in favor of newly arriving data). A *discard timer* permits removal of all packets from a queue that have been waiting for a certain time period. All those strategies can be extended to a *frame-based* buffer management, which drops all IP packets belonging to the same video frame if one of its IP packets was dropped.

Here, a proactive approach has been investigated that drops B-packets if a congestion situation is imminent (i.e., if the buffer occupancy exceeds a certain threshold δ). This will be referred to as *proactive B-dropping*. Packet- and frame-based strategies are studied here. Further, the removal of all involved interdependent frames in the queue, which become undecodable with the loss of a frame are also considered (i.e., *frame-based with interframe dependencies*) [Orl07].

The scenario used here for the performance evaluation comprised several (4 to 5) real UDP single-layer MPEG-4 video encoded flows, each with an average bit rate of 308 kbit/s. The data amount storable in the queue corresponds with 6 s. All UMTS mobile terminals move at 30 km/h, and both slow and fast fading conditions were modeled. A *Proportional Fair* (PF) scheduler [Cha04] was used at the MAC layer to assign resources to the different data streams. Because of other delays within the RAN, such as retransmissions, the

choice of discard timer value has to be smaller than that of the play-out buffer. Here, it was set to 5.35 s.

In order to evaluate the performance of the system, the (video) *Frame Error Rate* (FER) metric, which is reference-free, was used. FER describes the fraction of frames in error. If one IP packet in a video frame is lost, this frame and all other frames depending on this frame are considered to be frames in error. High or low FER values then stand for a bad or good perceived video quality, respectively.

In Fig. 2.27, a comparison of performance results for several proactive B-dropping schemes as a function of δ , the buffer occupancy threshold for proactive B-dropping, is presented. Both frame-based schemes show a very similar performance, though the consideration of interframe dependencies gives a slight advantage in certain ranges of δ at the cost of a much higher complexity. In contrast, the proactive packet-based scheme shows a worse performance. However, as a timer mechanism drops obsolete packets from the *Radio Network Controller* (RNC) input queue, the performance of the proactive packet-based scheme with timer greatly improves: in particular, the performance of such scheme is only weakly dependent on threshold δ . This allows for easy implementation without any data differentiation ($\delta = 1$). Additionally, it can be shown that proactive B-dropping, with optimal δ , results

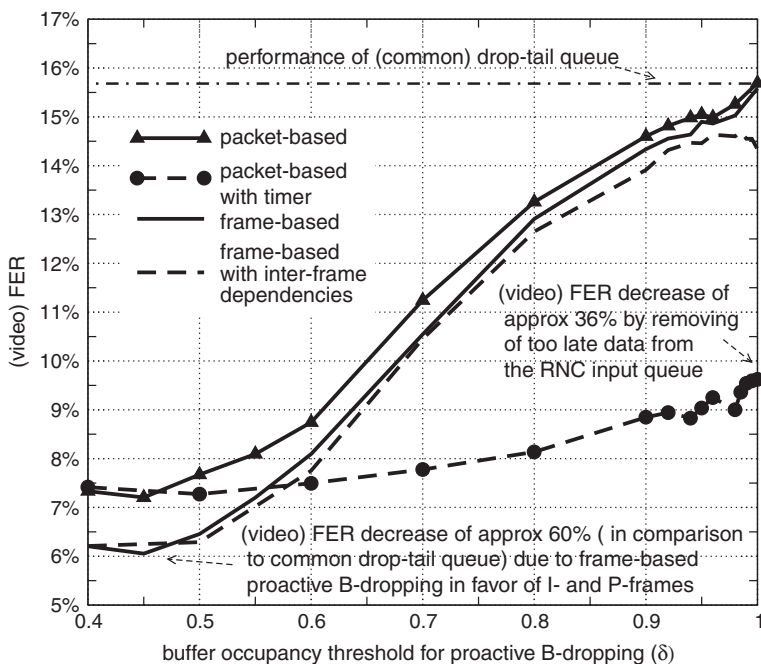


Fig. 2.27 Comparison of proactive B-dropping schemes

in (many) small interrupts, much better than a small number of (very) long video interrupts in the case without it.

2.6.4.4 GEO Satellite HSDPA Packet Scheduling

Satellite systems are a valid alternative to provide broadband communications to mobile and fixed users, complementing the coverage of terrestrial wireless and cellular systems; see also Section 2.5.2. The following study investigates packet scheduling aspects for *Satellite–Universal Mobile Telecommunication System* (S-UMTS) transmissions based on HSDPA (S-HSDPA) [ETSI00, TD(06)013, Gia07]. Here, it is considered a multi-spot-beam GEO bent-pipe satellite with all RAN functionalities corresponding with the network part located at the Node-B and the gateway on the earth.

Using the S-HSDPA multicode operation, several codes can be assigned to a UE, and several UEs can be scheduled in the same *Transmission Time Interval* (TTI). The UE reports the SINR experienced in terms of a *Channel Quality Indicator* (CQI) value with certain periodicity (in the studied scenario, 40 ms). The CQI value describes the modulation type (QPSK or 16QAM), the number of codes that can be used by the UE, and the corresponding maximum *Transport Block Size* (TBS) that guarantees a *Block Error Rate* (BLER) level below $BLER_{\text{threshold}}$.

The strict requirement $BLER_{\text{threshold}} = 0.01$ has been considered, as the *Round-Trip Propagation Delay* (RTPD) here is taken as 560 ms (GEO satellite case). Hence, the use of retransmissions to recover packet losses for real-time traffic is effectively prevented. Because the GEO RTPD is high, there will be a “misalignment” between the current SINR value at the UE and the CQI level that was used by the Node-B to transmit the transport block. To overcome this problem, the UE selects the CQI value by considering a suitable margin h [dB] on SINR ($h = 3.5$ dB in this study) [TD(06)013, Gia07].

For the simulation study here, distinct IP queues are used for different traffic flows according to the DiffServ approach. Two alternative packet scheduling techniques were considered, PF and *PF with Exponential Rule* (PF-ER) [And02, TD(05)045]. In both cases, scheduling decisions are taken at layer 3 according to layer 2 service parameters and PHY layer CQI information, thus resulting in an integrated cross-layer action. For the sake of comparison, results are also shown in the case of the *Earliest Deadline First* (EDF) scheduler (synonymous with EDD scheduler, used in Section 2.2.1) that bases its decisions only on the residual lifetime of queued IP packets.

In this study, video streaming and Web downloading traffic flows have been considered to be transmitted to UEs. *Standard Interchange Format* (SIF; 320×240 pixels) resolution (H.263 codec) is considered for video traffic, equivalent to 160 kbit/s (7.5 frames/s) per video stream. Web sources (2-state Markov-modulated Poisson arrival process of datagrams) have a mean bit-rate of 5.83 kbit/s. For each video packet (IP level), a lifetime of 160 ms (packet deadline) is used; after this time, it is cleared from the layer 3 buffer.

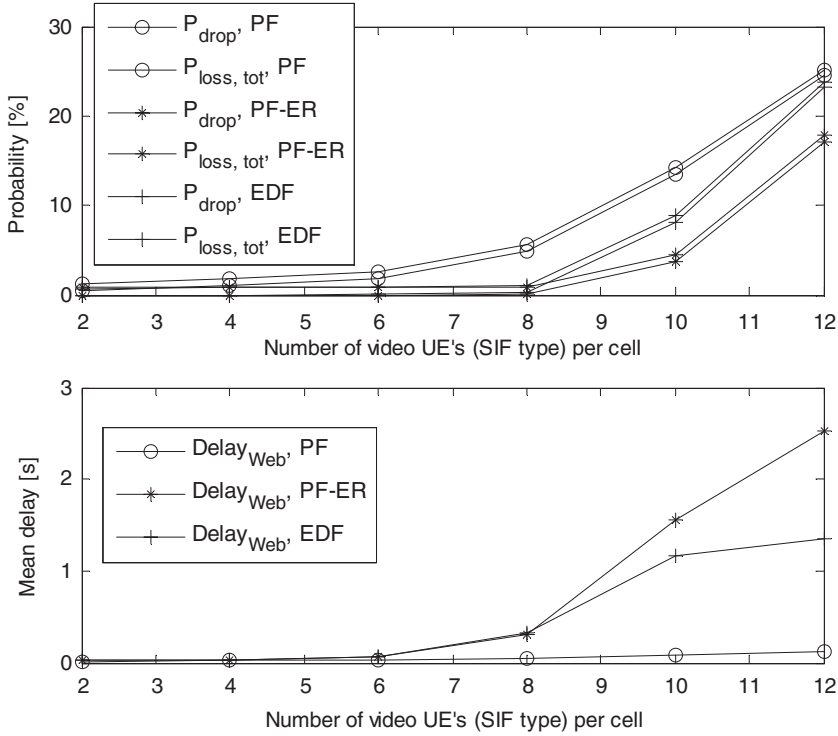


Fig. 2.28 Performance results, P_{drop} , $P_{loss,tot}$, and $Delay_{Web}$, for EDF, PF, and PF-ER scheduling schemes as a function of the number of video SIF UEs/cell for 50 Web UEs/cell

A preferential maximum delay for Web traffic of 500 ms has also been considered [TD(06)013, Gia07].

Figure 2.28 shows the layer 3 performance results as a function of the number of SIF video sources per cell with 50 Web traffic flows in terms of both the IP packet dropping probability due to deadline expiration for video traffic sources (P_{drop}) and the total IP packet loss probability ($P_{loss,tot}$). $P_{loss,tot}$ considers both P_{drop} and the losses introduced by the channel, due to the above-explained “misalignment” (in any case related to $BLER_{threshold} = 0.01$, due to the appropriate selection of the value of margin h).¹⁴ From these results, it may be noted that the P_{drop} sensitivity increases with the number of video UEs, as might be expected. The PF-ER scheme achieves the best performance for the video traffic management, and the PF technique is the best solution for the Web traffic performance in terms of the mean transmission delay for IP packet ($Delay_{Web}$). The reason is that the PF scheme selects the UE for transmissions

¹⁴ In the graph of Fig. 2.28, the $P_{loss,tot}$ curves closely follow the P_{drop} related curves because the P_{drop} term is the dominant one in the determination of $P_{loss,tot}$.

to distribute resources fairly among them, whereas the PF-ER technique bases its decisions also on deadlines, thus taking better into account the urgency of video packets over and above Web download packets.

With PF, the percentage of error-free frames obtained is 41.31% and for frames with invisible impairments (under the 36 dB threshold) it is 46.75%. Interestingly, with EDF the percentages are 4.88% and 9.59%, respectively. Hence, clearly, the EDF scheme does not permit a satisfactory video quality, whereas channel-aware schedulers represent a greatly improved solution.

2.7 Conclusions

This chapter focuses on resource management issues as crucial elements to support QoS in wireless systems. In particular, innovative MAC layer solutions, proposed in the COST 290 Action, have been presented, such as scheduling schemes, modified access protocols for current wireless standards, and CAC techniques. Congestion control has also been included in the studies, it being another important and crucial aspect for the massive access to, and usage growth of, the Internet through wireless systems. Different congestion control techniques have been presented with the aim of relating their performance to the behavior and the decisions taken at the MAC layer. Finally, innovative cross-layer design-based solutions have been presented. These solutions permit exploitation of the wireless system dynamics, which mainly arise because of the variability in the channel medium conditions and the use of a variety of transmission options on the physical layer, using interactions and relationships among the different protocol layers. Through a series of research-based experiments in various systems and scenarios, it has been shown that this innovative approach can enable capacity and QoS performance improvements for multimedia applications.

References

- [Aky01] I. Akyildiz, G. Morabito, S. Palazzo, TCP-Peach: A New Congestion Control Scheme for Satellite IP Networks, *IEEE/ACM Transactions on Networking*, Vol. 9, No. 3, pp. 307–321, June 2001.
- [Alo07] I. Alocci, G. Giambene, Y. Koucheryavy, Optimization of the Transport Layer Performance in a Wireless System Based on the IEEE 802.11e Standard, *International Symposium on Wireless Pervasive Computing 2007 (ISWPC2007)*, Puerto Rico (USA), February 5–7, 2007.
- [Alw96] A. Alwan, R. Bagrodia, N. Bambos, M. Gerla, L. Kleinrock, J. Short, J. Villasenor, *Adaptive Mobile Multimedia Networks*, IEEE Personal Communications, Vol. 3, No. 2, pp. 34–51, April 1996.
- [And02] M. Andrews, K. Kumaran, K. Ramanan, A. Stolyar, R. Vijayakumar, P. Whiting, CDMA Data QoS Scheduling on the Forward Link with Variable Channel Conditions, *Bell Labs Technical Memorandum*, April 2002.

- [Ath01] S. Athuraliya, V. H. Li, S. H. Low, Q. Yin, REM: Active Queue Management, *IEEE Network Magazine*, Vol. 15, No. 3, pp. 48–53, May/June 2001.
- [Bak85] J. E. Baker, Adaptive Selection Methods for Genetic Algorithms, 1st International Conference on Genetic Algorithms (ICGA), pp. 101–111, Lawrence Erlbaum Associates, Inc., Mahwah, NJ, USA, 1985.
- [Bal02] H. Balakrishnan, V. Padmanabhan, G. Fairhurst, M. Sooriyabandara, TCP Performance Implications of Network Path Asymmetry, RFC 3449, December 2002.
- [Bar02] G. Barriac, J. Holtzman, Introducing Delay Sensitivity into the Proportional Fair Algorithm for CDMA Downlink Scheduling, *IEEE International Symposium on Spread Spectrum Techniques and Applications 2002 (ISSTA2002)*, Parsippany, NJ, USA, Vol. 3, pp. 652–656, September 2002.
- [Bar03] C. A. Barnett, K. J. Ray Liu, Resource Efficient Multicast for 3G UMTS Wireless Networks, 58th IEEE Vehicular Technology Conference 2003 (VTC2003), 2003.
- [Bas06] S. A. Baset, H. Schulzrinne, An Analysis of the Skype Peer-to-Peer Internet Telephony Protocol, *IEEE INFOCOM 2006*, Barcelona, Spain, April 2006.
- [Bog06] K. Al-Begain, A. N. Dudin, V. V. Mushko, Novel Queuing Model for Multimedia Over Downlink in 3.5G Wireless Network, *Journal of Communications Software and Systems*, Vol. 2, No. 2, pp. 68–80, June 2006.
- [Bel06] B. Bellalta, M. Meo, M. Oliver, VoIP Call Admission Control in WLANs in presence of elastic traffic, *Journal of Communications Software and Systems*, Vol. 2, No. 4, December 2006.
- [Ber07] E. Bertran, M. S. O'Droma, P. L. Gilabert, G. Montoro, Performance Analysis of Power Amplifier Back-off Levels in UWB Transmitters, *IEEE Transactions on Consumer Electronics*, Vol. 53, No. 4, pp. 1309–1313, November 2007.
- [Blo07] S. Romaszko, C. Blondia, Neighbor and Interference-Aware MAC Protocol for Wireless ad hoc Networks, *IST Mobile and Wireless Communication Summit 2007*, Hungary, Budapest, July 2007.
- [Bog07] G. Boggia, P. Camarda, L. A. Grieco, S. Mascolo, Feedback-based Control for Providing Real-time Services with the 802.11e MAC, *IEEE/ACM Transactions on Networking*, Vol. 15, No. 2, pp. 323–333, April 2007.
- [Bog06] G. Boggia, P. Camarda, L. A. Grieco, S. Mascolo, Energy Efficient Feedback-based Scheduler for Delay Guarantees in IEEE 802.11e, *Networks Computer Communications*, special issue, Vol. 29, No. 3–4, pp. 2680–2692, August 2006.
- [Bon05] T. Bonald, S. Borst, A. Proutière, Inter-Cell Scheduling in Wireless Data Networks, *European Wireless 2005 (EW 2005)*, Nicosia, Cyprus, April 2005.
- [Bu06] T. Bu, Y. Liu, D. Towsley, On the TCP-Friendliness of VoIP Traffic, *IEEE INFOCOM 2006*, pp. 1–12, Barcelona, Spain, April 2006.
- [Cas05] C. Casetti, C.-F. Chiasserini, L. Merello, G. Olmo, Supporting Multimedia Traffic in 802.11e WLANs, 61st IEEE Vehicular Technology Conference 2005 (VTC 2005-Spring), Vol. 4, pp. 2340–2344, 30 May – 1 June 2005.
- [Cha04] N. Challa, H. Cam, Cost-Aware Downlink Scheduling of Shared Channels for Cellular Networks with Relays, *IEEE International Performance of Computers and Communication Conference 2004 (IPCCC2004)*, pp. 793–798, Phoenix, AZ, USA, April 15–17, 2004.
- [Cha05] S. Chang, A. Vetro, Video Adaptation: Concepts, Technologies and Open Issues, *Proceedings of IEEE*, Vol. 93, No. 1, pp. 148–158, January 2005.
- [Che02] K. Chen, S. Shan, K. Nahrstedt, Cross-layer Design for Accessibility in Mobile ad hoc Networks, *International Journal on Wireless Personal Communications*, Vol. 21, No. 1, pp. 49–76, April 2002.
- [Che05] L. Chen, S. H. Low, J. C. Doyle, Joint Congestion Control and Media Access Control Design for Ad Hoc Wireless Networks, *IEEE INFOCOM 2005*, 2005.
- [Che06a] K. Chen, C. Huang, P. Huang, C. Lei, Quantifying Skype User Satisfaction, *Special Interest Group on Data Communication 2006 (SIGCOMM2006)*, pp. 339–410, Pisa, Italy, September 2006.

- [Che06b] J. J. Chen, L. Lee, Y. C. Tseng, Integrating SIP and IEEE 802.11e to Support Handoff and Multi-grade QoS for VoIP Applications, ACM International Workshop on QoS and Security for Wireless and Mobile Networks 2006 (Q2SWINET2006), pp. 67–74, Torremolinos, Malaga, Spain, October 2–6, 2006.
- [Chr06] C. Chrysostomou, A. Pitsillides, Fuzzy Logic Congestion Control in TCP/IP Tandem Networks, IEEE Symposium on Computers and Communications 2006 (ISCC2006), pp. 123–129, July 2006.
- [Cic07] L. De Cicco, S. Mascolo, V. Palmisano An Experimental Investigation of the Congestion Control Used by Skype VoIP, 5th International Conference on Wired/Wireless Internet Communications 2007 (WWIC 2007), pp. 153–164, Coimbra, Portugal, May 2007.
- [Col99] G. Colombo, L. Lenzini, E. Mingozzi, B. Cornaglia, R. Santaniello, Performance Evaluation of PRADOS: a Scheduling Algorithm for Traffic Integration in Wireless ATM Networks, ACM/IEEE International Conference on Mobile Computing and Networking 1999 (MOBICOM1999), Seattle, WA, pp. 143–150, August 1999.
- [Com84] R. A. Comroe, D. J. Costello Jr., ARQ Schemes for Data Transmission in Mobile Radio Systems, IEEE Journal of Selected Areas in Communications, Vol. 2, No. 4, pp. 472–481, July 1984.
- [Cos05] X. P. Costa, D. C. Mur, T. Sashihara, Analysis of the Integration of IEEE 802.11e Capabilities in Battery Limited Mobile Devices, IEEE Wireless Communications, Vol. 12, No. 6, pp. 26–32, December 2005.
- [COS290] Web site of the COST 290 Action with URL (date of access January 2008): <http://www.cost290.org/>.
- [Dav02] B. Davie, A. Charny, J. C. R. Bennet, K. Benson, J. Y. Le Boudec, W. Courtney, S. Davari, V. Firoiu, D. Stiliadis, An Expedited Forwarding PHB (Per-Hop Behavior), RFC 3246, March 2002.
- [Del05] F. M. Delicado, P. Cuenca, L. Orozco-Barbosa, Multiservice Communications over TDMA/TDD wireless LANs, 3rd International Conference on Wired/Wireless Internet Communications 2005 (WWIC 2005), Xanthi, Greece, 2005.
- [Del06] F. M. Delicado, P. Cuenca, L. Orozco-Barbosa, QoS Mechanisms for Multimedia Communications over TDMA/TDD WLANs, Computer Communications Journal, Vol. 29, No. 13–14, pp. 2721–2735, August 2006.
- [Dem89] A. Demers, S. Keshav, S. Shenker, Analysis and Simulation of a Fair Queueing Algorithm, Special Interest Group on Data Communication 1989 (SIGCOMM1989), Austin, TX, pp. 1–12, 1989.
- [Duk05] N. Dukkupati, M. Kobayashi, R. Zhang-Shen, N. McKeown, Processor Sharing Flows in the Internet, International Workshop on Quality of Service 2005 (IWQoS2005), pp. 267–281, June 2005.
- [Ela05] S. E. Elayoubi, T. Chahed. Admission Control in the Downlink of WCDMA/UMTS. Springer-Verlag Berlin, Heidelberg, pp. 136–151, 2005.
- [Ela04] S. E. Elayoubi, T. Chahed, G. Hébuterne, Connection Admission Control in UMTS in the Presence of Shared Channels, Computer Communications, Vol. 27, No. 11, June 2004.
- [Ele95] A. Eleftheriadis, D. Anastassiou, Meeting Arbitrary QoS Constraints Using Dynamic Rate Shaping of Coded Digital Video, International Workshop on Networking and Operating System Support for Digital Audio and Video 1995 (NOSSDAV1995), pp. 89–100, Durham, NH, USA, April 19–21, 1995.
- [Ent07] J. T. Entrambasaguas, M. C. Aguayo-Torres, G. Gomez, J. F. Paris, Multiuser Capacity and Fairness Evaluation of Channel/QoS-Aware Multiplexing Algorithms, IEEE Network, Vol. 21, No. 3, pp. 24–30, May–June 2007.
- [ETSI00] ETSI, Part 1 to 4 of Satellite Earth Stations and Systems (SES); Satellite Component of UMTS/IMT2000; G-family, TS 101 851.
- [Fel06] F. Feller, M. Necker, Comparison of Opportunistic Scheduling Algorithms for HSDPA Networks, 12th EUNICE Summer School, Stuttgart, Germany, September 2006.

- [Flo91] S. Floyd, V. Jacobson, Connections with Multiple Congested Gateways in Packet-switched Networks, *Computer Communications Review*, Vol. 21, No. 5, pp. 30–47, August 1991.
- [Flo93] S. Floyd, V. Jacobson, Random Early Detection Gateways for Congestion Avoidance, *IEEE/ACM Transactions on Networking*, Vol. 1, No. 4, pp. 397–413, August 1993.
- [Flo99] S. Floyd, T. Henderson, The NewReno Modification to TCP's Fast Recovery Algorithm, RFC 2582, 1999.
- [Flo01] S. Floyd, R. Gummadi, S. Shenker, Adaptive RED: An Algorithm for Increasing the Robustness of RED's Active Queue Management, Technical report, ICSI, 2001.
- [Flo03] S. Floyd, M. Handley, J. Padhye, J. Widmer, TCP Friendly Rate Control (TFRC): Protocol Specification, RFC 3448, January 2003.
- [Flo04] S. Floyd, J. Kempf, IAB Concerns Regarding Congestion Control for Voice Traffic in the Internet, RFC 3714, March 2004.
- [Fri04] V. Friderikos, L. Wang, A. H. Aghvami, TCP-aware Power and Rate Adaptation in DS/CDMA Networks, *IEE Proceedings: Communications*, Vol. 151, No. 6, pp. 581–588, December 2004.
- [Flo07a] S. Floyd, M. Allman, A. Jain, P. Sarolahti, Quick-Start for TCP and IP, IETF RFC 4782 (experimental), January 2007.
- [Flo07b] S. Floyd, Metrics for the Evaluation of Congestion Control Mechanisms, Internet draft – IETF, Expires: August 2007.
- [Gia07] G. Giambene, S. Giannetti, C. Párraga Niebla, M. Ries, A. Sali, Traffic Management in HSDPA via GEO Satellite, *Space Communications Journal*, Vol. 21, No. 1–2, pp. 37–54, 2007/2008.
- [Gia06] G. Giambene, S. Kota, Cross-layer Protocol Optimization for Satellite Communications Networks: A Survey, *International Journal of Satellite Communications and Networking*, Vol. 24, No. 5, pp. 323–341, September–October 2006.
- [Gom07] G. Gómez, J. Poncea González, M. C. Aguayo-Torres, J. F. Paris, J. T. Entrambasaguas, QoS Modeling for Performance Evaluation over Evolved 3G Networks, *ACM International Workshop on QoS and Security for Wireless and Mobile Networks 2007 (Q2SWINET2007)*, pp. 148–151, Chania, Crete Island, Greece, October 26, 2007.
- [Hei99] J. Heinanen, F. Baker, W. Weiss, J. Wroclawski, Assured Forwarding PHB Group, Internet standard RFC 2597, June 1999.
- [Hem99] M. Hemy, U. Hengartner, P. Steenkiste, T. Gross, MPEG System Streams in Best-Effort Networks, *Packet Video Workshop 1999 (PV1999)*, May 1999.
- [Hen99] T. R. Henderson, R. H. Katz, Transport Protocols for Internet-Compatible Satellite Networks, *IEEE Journal on Selected Areas in Communications*, Vol. 17, No. 2, pp. 345–359, February 1999.
- [Heu03] M. Heusse, F. Rousseau, G. Berger-Sabbatel, A. Duda, Performance Anomaly of 802.11b, *IEEE INFOCOM2003*, pp. 836–842, San Francisco, CA, USA, April 1–3, 2003.
- [Hol01] H. Holma, A. Toskala. WCDMA for UMTS. 3rd Edition, John Wiley & Sons Ltd, Chichester, England, 2001.
- [Hol02] C. V. Hollot, V. Misra, D. Towsley, W.-B. Gong, Analysis and Design of Controllers for AQM Routers Supporting TCP Flows, *IEEE Transactions on Automatic Control*, Vol. 47, No. 6, pp. 945–959, June 2002.
- [Hos04] E. Hossain, V. K. Bhargava, Cross-layer Performance in Cellular WCDMA/3G Networks: Modelling and Analysis, *IEEE International Symposium on Personal, Indoor and Mobile Radio Communications 2004 (PIMRC2004)*, Vol. 1, pp. 437–443, Barcelona, Spain, September 2004.
- [IEEE05] IEEE 802 Committee of the IEEE Computer Society, IEEE P802.11e Amendment to IEEE Std 802.11, Part 11: Wireless LAN Medium Access Control (MAC) and Physical Layer (PHY) Specifications: Medium Access Control (MAC) Quality of Service (QoS) Enhancements, November 2005.

- [IEEE99] LAN MAN Standards Committee of the IEEE Computer Society, ANSI/IEEE Std 802.11, Part 11: Wireless LAN Medium Access Control (MAC) and Physical Layer (PHY) Specifications, 1999 Edition, 1999.
- [IEEE9] The Institute of Electrical and Electronic Engineers, IEEE Computer Society LAN MAN Standards Committee: Wireless LAN Medium Access Control (MAC) and Physical Layer (PHY) Specifications, ANSI/IEEE Std. 802.11, ANSI/IEEE Std. 802.11, 1999 Edition, 1999.
- [ISO99] Information technology – Generic coding of audio-visual objects- Part 2: Visual, ISO Std. ISO/IEC 14486-2 PDAM1, 1999.
- [ITU05] Advanced Video Coding for Generic Audiovisual Services, ITU-T Recommendation H.264, 2005.
- [ITU92] Coding of Speech at 16 kbit/s Using Low-delay Code Excited Linear Prediction, Std. ITU-T Recommendation G.728, September 1992.
- [Jac55] J. Jackson, Scheduling a Production Line to Minimize Maximum Tardiness, Research Report No. 43, Management Science Research Project, University of California, Los Angeles, 1955.
- [Jai03] K. Jain, J. Padhye, V. N. Padmanabhan, L. Qiu, Impact of Interference on Multi-hop Wireless Network Performance, 9th ACM/IEEE International Conference on Mobile Computing and Networking 2003 (MOBICOM2003), San Diego, CA, USA, pp. 66–80, 2003.
- [Jal00] A. Jalali, R. P. R. Pankaj, Data Throughput of CDMA-HDR a High Efficiency-high Data Rate Personal Communication Wireless System, Vehicular Technology Conference 2000 (VTC2000-Spring), Tokyo, Japan, May 2000.
- [Kar00] A. Karam, F. Tobagi, On the Traffic and Service Classes in the Internet, IEEE GLOBECOM2000, San Francisco, CA, USA, 2000.
- [Kat02] D. Katabi, M. Handley, C. Rohrs, Internet Congestion Control for High-Bandwidth-Delay Products, Special Interest Group on Data Communication 2002 (SIGCOMM2002), Pittsburgh, PA, August 2002.
- [Kle01] A. Klemm, C. Lindemann, M. Lohmann, Traffic Modeling and Characterization for UMTS Networks, IEEE GLOBECOM2001, Internet Performance Symposium, San Antonio, TX, November 2001.
- [Kha04] A. K. F. Khattab, K. M. F. Elsayed, Channel-quality Dependent Earliest Deadline Due Fair Scheduling Schemes for Wireless Multimedia Networks, ACM International Symposium on Modeling, Analysis and Simulation of Wireless and Mobile Systems 2004 (MSWiM 2004), Venice, Italy, October 2004.
- [Kim05] N. Kim, H. Yoon, Wireless Packet Fair Queuing Algorithms with Link Level Retransmission, Computer Communications, Vol. 28, No. 7, pp. 713–725, May 2005.
- [Kol03] T. Kolding, Link and System Performance Aspects of Proportional Fair Scheduling in WCDMA/HSDPA, Vehicular Technology Conference 2003 (VTC 2003-Fall), pp. 1717–1722, October 2003.
- [Kse06] A. Ksentini, M. Naimi, A. Gueroui, Toward an Improvement of H.264 Video Transmission over 802.11e through a Cross-Layer Architecture, IEEE Communication Magazine, Vol. 44, No.1, pp. 107–114, January 2006.
- [Kun04] S. Kunniyur, R. Srikant, An Adaptive Virtual Queue (AVQ) Algorithm for Active Queue Management, IEEE/ACM Transactions on Networking, Vol. 12, No. 2, pp. 286–299, April 2004.
- [Kwo04] Y. Kwon, Y. Fang, H. Latchman, Design of MAC Protocols with Fast Collision Resolution for Wireless Local Area Networks, IEEE Transactions on Wireless Communications, Vol. 3, No. 3, pp. 793–807, May 2004.
- [Lak97] T. V. Lakshman, U. Madhow, The Performance of TCP/IP for Networks with High Bandwidth-Delay Products and Random Loss. IEEE/ACM Transactions on Networking, Vol. 5, No. 3, pp. 336–350, June 1997.

- [Lei05] D. J. Leith, P. Clifford, TCP Fairness in 802.11e WLANs, International Conference on Wireless Networks 2005 (ICWN2005), 2005.
- [Les07] M. Lestas, A. Pitsillides, P. Ioannou, G. Hadjipollas, ACP: A Congestion Control Protocol with Learning Capability, Computer Networks, Vol. 51, No. 13, pp. 3773–3798, September 2007.
- [Liu06] S. Liu, J. Virtamo, Inter-Cell Coordination with Inhomogeneous Traffic Distribution, 2nd Conference on Next Generation Internet Design and Engineering 2006 (NGI2006), València, Spain, April 2006.
- [Liu73] C. L. Liu, J. W. Layland, Scheduling Algorithms for Multiprogramming in a Hard-real-time Environment, Journal of the ACM, Vol. 20, No. 1, pp. 46–61, January 1973.
- [Lou07] P. Loureiro, S. Mascolo, E. Monteiro, Open Box Protocol (OBP), High Performance Computing and Communications 2007 (HPCC 2007), Springer LNCS, Houston, TX, USA, pp. 496–507, September 2007.
- [Low02] S. H. Low, F. Paganini, J. Wang, S. Adlakha, J. C. Doyle, Dynamics of TCP/RED and a Scalable Control, IEEE INFOCOM2002, Vol. 1, pp. 23–27, June 2002.
- [Low05] S. H. Low, L. L. H. Andrew, B. P. Wydrowski, Understanding XCP: Equilibrium and fairness, IEEE INFOCOM2005, Vol. 2, pp. 1025–1036, March 2005.
- [Lun05] T. Lundberg, P. de Bruin, S. Bruhn, S. Hakansson, S. Craig, Adaptive Thresholds for AMR Codec Mode Selection, Vehicular Technology Conference 2005 (VTC 2005-Spring), Vol. 4, pp. 2325–2329, Dallas, TX, USA, September 25–28, 2005.
- [Mal04] M. Malli, Q. Ni, T. Turletti, C. Barakat, Adaptive Fair Channel Allocation for QoS Enhancement in IEEE 802.11 Wireless LANs, IEEE International Conference on Communications 2004 (ICC2004), Paris, France, June 2004.
- [Mam05] L. Mamas, V. Tsaoussidis, A new Approach to Service Differentiation: Non-Congestive Queuing, ICST First International Workshop on Convergence of Heterogeneous Wireless Networks 2005 (CONWIN2005), Budapest, Hungary, July 2005.
- [Mam07] L. Mamas, V. Tsaoussidis, Differentiating Services for Sensor Internetworking, the IFIP Fifth Annual Mediterranean Ad Hoc Networking Workshop 2007 (Med-Hoc-Net 2007), Corfu, Greece, June 2007.
- [Man03] S. Mangold, S. Choi, G. R. Hiertz, O. Klein, B. Walke, Analysis of IEEE 802.11e for QoS Support In Wireless LANs, IEEE Wireless Communication Magazine, Vol. 10, No. 6, pp. 40–50, December 2003.
- [Mas99] S. Mascolo, Congestion Control in High-Speed Communication Networks Using the Smith Principle, Automatica, Vol. 35, No. 12, pp. 1921–1935, December 1999. Special Issue on Control methods for communication networks.
- [Mas01] S. Mascolo, C. Casetti, M. Gerla, M. Sanadidi, R. Wang, TCP Westwood: Bandwidth Estimation for Enhanced Transport over Wireless Links, ACM/IEEE International Conference on Mobile Computing and Networking 2001 (MOBICOM2001), Rome, Italy, July 2001.
- [Min00] T. Minn, K.-Y. Siu, Dynamic Assignment of Orthogonal Variable-Spreading-Factor Codes in W-CDMA, Journal of Selected Areas in Communications, Vol. 18, No. 8, pp. 1429–1440, August 2000.
- [Nec06] M. C. Necker, A Comparison of Scheduling Mechanisms for Service Class Differentiation in HSDPA Networks, International Journal of Electronics and Communications (AEÜ), Vol. 60, No. 2, pp. 136–141, February 2006.
- [Nec07a] M. C. Necker, Integrated Scheduling and Interference Coordination in Cellular OFDMA Networks, IEEE International Conference on Broadband Communications, Networks and Systems 2007 (BROADNETS 2007), Raleigh, NC, USA, September 2007.
- [Nec07b] M. C. Necker, Coordinated Fractional Frequency Reuse, 9th ACM/IEEE International Symposium on Modeling, Analysis and Simulation of Wireless and Mobile Systems 2007 (MSWiM2007), Chania, Crete Island, October 2007.

- [Neo06] M. Neophytou, A. Pitsillides, Hybrid CAC for MBMS-Enabled 3G UMTS Networks, 17th Annual IEEE International Symposium on Personal, Indoor and Mobile Radio Communications 2006 (PIMRC2006), Helsinki, Finland, September 2006.
- [Nie06] I. Niemegeer, Layerless Networking?, 8th Strategic Workshop 2006 (SW2006), Mykonos, Greece, June 2–4, 2006.
- [ns207] Ns-2. Network simulator, available at the URL (date of access, January 2008): <http://www.isi.edu/nsnam/ns>.
- [ODR05] M. O'Droma, E. Bertran, J. Portilla, N. Mgebrishvili, S. Donati Guerrieri, G. Montoro, T. J. Brazil, G. Magerl, On Linearisation of Microwave Transmitter Solid State Power Amplifiers, Wiley International Journal of RF and Microwave Computer-Aided Engineering (RFMiCAE), Vol. 15, No. 5, pp. 491–505, September 2005.
- [OPNET04] Opnet Technologies Inc. OPNET Modeler 10.0, 1987–2004, available at the URL (date of access, January 2008): <http://www.opnet.com>.
- [Orl07] Z. Orlov, M. C. Necker, Enhancement of Video Streaming QoS with Active Buffer Management in Wireless Environments, European Wireless 2007 (EW2007), Paris, France, April 1–4, 2007.
- [Pad98] J. Padhye, V. Firoiu, D. Towsley, J. Kurose, Modeling TCP Throughput: A Simple Model and Its Empirical Validation, ACM SIGCOMM Computer Communication Review, Vol. 28, No. 4, pp. 303–314, October 1998.
- [Par93] A. K. Parekh, R. G. Gallager, A Generalized Processor Sharing Approach to Flow Control in Integrated Services Networks: the Single-node Case, IEEE/ACM Transactions on Networking, Vol. 1, No. 3, pp. 344–357, June 1993.
- [Pri04] J. Price, T. Javidi, Cross-layer (MAC and Transport) Optimal Rate Assignment in CDMA-based Wireless Broadband Networks, 38th Asilomar Conference on Signals, Systems and Computers 2004 (ACSSC2004), Vol. 1, pp. 1044–1048, Pacific Grove, CA, USA, November 2004.
- [R1-051051] 3GPP TSG RAN WG1#42 R1-051051: Standard Aspects of Interference Coordination in EUTRA, 2005.
- [Ram89] R. Ramaswami, K. K. Parhi, Distributed Scheduling of Broadcasts in a Radio Network, IEEE INFOCOM 1989, Ottawa, ON, Canada, pp. 497–504, 1989.
- [Rom06] S. Romaszko, C. Blondia, Neighbour-aware, Collision Avoidance MAC Protocol (NCMac) for Mobile ad hoc Networks, International Symposium on Wireless Communication Systems 2006 (ISWCS2006), Spain, Valencia, September 2006.
- [Rom07] S. Romaszko, C. Blondia, Bounds Selection-dynamic Reset Protocol for Wireless ad hoc LANs, IEEE Wireless Communications & Networking Conference 2007 (WCNC2007), Hong-Kong, March, 2007.
- [Sar07] P. Sarolahti, M. Allman, S. Floyd, Determining an Appropriate Sending Rate over an Underutilized Network Path, Computer Networks, Vol. 51, No. 7, pp. 1815–1832, May 2007.
- [Sch07] M. Scharf, Performance Analysis of the Quick-Start TCP Extension, 4th IEEE International Conference on Broadband Communications, Networks and Systems 2007 (BROADNETS 2007), Raleigh, NC, USA, September 2007.
- [Sha01] S. Shakkottai, A. Stolyar, Scheduling Algorithms for a Mixture of Real-time and Non-real-time data in HDR, 17th International Teletraffic Congress 2001 (ITC2001), Salvador da Bahia, Brazil, September 2001.
- [Sim07] A. Simonsson, Frequency Reuse and Intercell Interference Coordination in E-UTRA, 65th IEEE Vehicular Technology Conference 2007 (VTC2007-Spring), pp. 3091–3095, Dublin, Ireland, April 2007.
- [Sri05] V. Srivastava, M. Motani, Cross-layer Design: a Survey and the Road Ahead, IEEE Communications Magazine, Vol. 43, No. 12, pp. 112–119, December 2005.
- [Ste97] W. Stevens, TCP Slow Start, Congestion Avoidance, Fast Retransmit, and Fast Recovery Algorithms, RFC 2001, January 1997.

- [Sud01] P. Sudame, B. Badrinath, On Providing Support for Protocol Adaptation in Mobile Wireless Networks, *Mobile Networks and Applications*, Vol. 6, No. 1, pp. 43–55, January/February 2001.
- [Sun06] J. Yuan Sun, L. Zhao, A. Anpalagan, Cross-Layer Design and Analysis of Downlink Communications in Cellular CDMA Systems, *EURASIP Journal on Wireless Communications and Networking*, Vol. 2006, No. 2, pp. 1–23, April 2006.
- [Syl06] S. Romaszko, C. Blondia, Neighbour and Energy-Aware Contention Avoidance MAC Protocol for Wireless Ad Hoc Networks, *IEEE International Conference on Wireless and Mobile Computing, Networking and Communications 2006 (WiMob2006)*, Montreal, Canada, June 2006.
- [Tse01] Y.-C. Tseng, C.-M. Chao, S.-L. Wu, Code Placement and Replacement Strategies for Wideband CDMA OVFS Code Tree Management, *IEEE GLOBECOM2001*, Vol. 1, pp. 562–566, 2001.
- [Vil05] R. Vilmann, C. Bettstetter, C. Hartmann, On the Impact of Beamforming on Interference in Wireless Mesh Networks, *IEEE Workshop on Wireless Mesh Networks 2005 (WiMesh2005)*, Santa Clara, CA, USA, September 2005.
- [Vil06a] J. Villalón, P. Cuenca, L. Orozco-Barbosa, B-EDCA: A New IEEE 802.11e-based QoS Protocol for Multimedia Wireless Communications, *IFIP Networking Conference 2006*, Coimbra, Portugal, May 2006.
- [Vil06b] J. Villalón, P. Cuenca, L. Orozco-Barbosa, A Novel IEEE 802.11e-based QoS Protocol for Voice Communications over WLANs, *4th International Conference on Wired/Wireless Internet Communications 2006 (WWIC2006)*, Bern, Switzerland, May 2006.
- [Vlc07] VideoLan – VLC Media Player, available at the URL (date of access, January 2008): <http://www.videolan.org>.
- [Vlo05] N. Vlotomas, J. Antoniou, G. Hadjipollas, V. Vassiliou, A. Pitsillides, Power Control for Efficient Multicasting in IP-based 3G and Beyond Mobile Networks, *11th European Wireless Conference (EW2005)*, Nicosia, Cyprus, April 2005.
- [Wu01] D. Wu, Y. T. Hou, Y. Q. Zhang, Scalable Video Coding and Transport over Broad-Band Wireless Networks, *Proc. IEEE*, Vol. 89, No. 1, pp. 6–20, January 2001.
- [Wu05] Y. Wu, Q. Zhang, W. Zhu, S.-Y. Kung, Bounding the Power Rate Function of Wireless ad hoc Networks, *IEEE INFOCOM 2005*, Miami, FL, USA, pp.584–595, March 2005.
- [Yan00] Y. R. Yang, S. S. Lam, General AIMD Congestion Control, *8th International Conference on Network Protocols 2000 (ICNP2000)*, Osaka, Japan, November 2000.
- [Zha95] H. Zhang, Service Disciplines for Guaranteed Performance Service in Packet-Switching Networks, *Proc. IEEE*, Vol. 83, No. 10, pp. 1374–1396, October 1995.
- [Zim80] H. Zimmermann, OSI Reference Model – The ISO Model of Architecture for Open Systems Interconnection, *IEEE Transactions on Communications*, Vol. 28, No. 4, pp.425–432, April 1980.

The COST 290 documents can be downloaded from the link www.cost290.org. The COST 290 documents referenced in this chapter are listed below:

- [STSM(06)002] I. Alocci, Optimization of the Transport Layer Performance in a Wireless System based on the IEEE 802.11e Standard.
- [TD(05)011] M. Cintează, T. Rădulescu, I. Marghescu, Orthogonal-Variable-Spreading-Factor Code Allocation Strategy using Genetic Algorithms.
- [TD(05)013] L. Mamatas, V. Tsaoussidis, A New Approach to Service Differentiation: Non-Congestive Queuing.
- [TD(05)032] G. Boggia, P. Camarda, F. A. Favia, L. A. Grieco, S. Mascolo, Energy Efficient Feedback-based Scheduler for Delay Guarantees in IEEE 802.11e Networks.

- [TD(05)045] J. T. Entrambasaguas, M. C. Aguayo-Torres, G. Gomez, J. F. Paris, Multiuser Capacity and Fairness Evaluation of Channel/QoS Aware Multiplexing Algorithms.
- [TD(05)048] K. Al-Begain, A. N. Dudin, V. Mushko, Analytical Model for Multimedia Provisioning over Downlink in 3.5G Wireless Networks.
- [TD(06)012] B. Bellalta, M. Meo, M. Oliver, Call Admission Control in IEEE 802.11e EDCA-based WLANs (Initial Steps).
- [TD(06)013] G. Giambene, S. Giannetti, V. Y. H. Kueh, C. Párraga Niebla, HSDPA and MBMS Transmissions via S-UMTS.
- [TD(06)015] M. C. Necker, Scheduling Mechanisms for QoS Differentiation in HSDPA Networks.
- [TD(06)028] S. Romaszko, C. Blondia, Neighbour and Energy-Aware Contention Avoidance MAC Protocol for Wireless Ad Hoc Networks.
- [TD(06)034] P. Papadimitriou, V. Tsaoussidis, Evaluating TCP Mechanisms for Real-Time Streaming over Satellite Links.
- [TD(06)038] J. Villalón, P. Cuenca, L. Orozco-Barbosa, A Novel IEEE 802.11e-based QoS Protocol for Voice Communications over WLANs.
- [TD(06)046] M. Necker, Global Interference Coordination in 802.16e Networks.
- [TD(07)012] Z. Orlov, Improvement of Video Streaming QoS by Application-aware Queue Management in UMTS/HSDPA Networks.
- [TD(07)013] M. Scharf, Quick-Start TCP: Performance Evaluation and Open Issues.
- [TD(07)014] G. Gómez, J. Poncela-González, M. C. Aguayo-Torres, J. F. Paris, J. T. Entrambasaguas, Impact of Transport Protocols in Application Throughput for Wireless-wired Networks.
- [TD(07)015] C. Cano, B. Bellalta, Flow-Level Evaluation of Call Admission Control Schemes in WMM WLANs.
- [TD(07)018] A. Sfairopoulou, C. Macián, B. Bellalta, Dynamic Measurement-based Codec Selection for VoIP in Multirate 802.11.
- [TD(07)021] G. Lazar, V. Dobrota, T. Blaga, Cross-Layer Architecture for H.264 Video Streaming in Heterogeneous DiffServ Networks.
- [TD(07)028] S. Romaszko, C. Blondia, Controlled Contention-based Access to the Medium in ad hoc WLANs.
- [TD(07)037] G. Lazar, V. Dobrota, T. Blaga, Performance of Wireless IEEE 802.11e-Based Devices with Multiple Hardware Queues.
- [TD(07)042] A. Sfairopoulou, B. Bellalta, C. Macián, Joint Admission Control and Adaptive VoIP Codec Selection Policies in Multi-rate WLANs.
- [TD(07)043] D. Moltchanov, Cross-layer Performance Control of Wireless Channels.
- [TD(07)050] S. Mascolo, An Experimental Investigation of the Congestion Control Used by Skype.

Traffic and QoS Management in Wireless Multimedia
Networks

COST 290 Final Report

Koucheryavy, Y.; Giambene, G.; Staehle, D.;

Barcelo-Arroyo, F.; Braun, T.; Siris, V. (Eds.)

2009, XVI, 312 p. 111 illus., Hardcover

ISBN: 978-0-387-85572-1